



BELL & HOWELL SCHOOLS

DeVRY INSTITUTE OF TECHNOLOGY



 **BELL & HOWELL SCHOOLS**

DEVRY INSTITUTE OF TECHNOLOGY

BASIC ELECTRONIC CIRCUIT APPLICATIONS

COURSE 252



BASIC ELECTRONIC CIRCUIT APPLICATIONS

CONTENTS

TRANSDUCERS	Lesson 1082	
ELECTRONIC SYSTEMS	Lesson 1085	
VACUUM TUBE HIGH FREQUENCY AMPLIFIERS	Lesson 1088	
TRANSISTOR PUSH-PULL AND HIGH FREQUENCY AMPLIFIERS....	Lesson 1091	
VACUUM TUBE OSCILLATORS	Lesson 1094	37
TRANSISTOR OSCILLATORS	Lesson 1097	27
FIELD EFFECT TRANSISTORS	Lesson 1100	24
SEMICONDUCTOR CONTROL DEVICES	Lesson 1103	30
MOTORS AND GENERATORS	Lesson 1106	42
PRINTED CIRCUIT TECHNIQUES	Lesson 1109	29
INTEGRATED CIRCUITS	Lesson 1112	34
ELECTROMECHANICAL DEVICES	Lesson 1115	36
UNIT EXAMINATION	Course 252	

BASIC ELECTRONIC CIRCUIT APPLICATIONS

CONTENTS

1. Introduction	1
2. Basic Electronic Components	2
3. Basic Electronic Circuits	3
4. Basic Electronic Applications	4
5. Basic Electronic Troubleshooting	5
6. Basic Electronic Safety	6
7. Basic Electronic Maintenance	7
8. Basic Electronic Repair	8
9. Basic Electronic Replacement	9
10. Basic Electronic Testing	10
11. Basic Electronic Calibration	11
12. Basic Electronic Adjustment	12
13. Basic Electronic Assembly	13
14. Basic Electronic Disassembly	14
15. Basic Electronic Packaging	15
16. Basic Electronic Labeling	16
17. Basic Electronic Marking	17
18. Basic Electronic Identification	18
19. Basic Electronic Documentation	19
20. Basic Electronic Record Keeping	20

TRANSDUCERS

1082

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





Transducers are mechanical devices which are in such widespread use that their importance in daily life is easily overlooked. We take for granted the recreational uses for transducers in radio, television, and recording equipment. Many handicapped persons, however, rely upon the use of cardiac pacemakers, hearing aids, electronic canes, and other transducers which help them live more useful lives.

Courtesy Zenith Sales Corp.

CONTENTS

Energy Converters	Page 4
Variable Resistance Transducers	Page 6
Variable Reluctance Transducers	Page 8
Variable Capacitance Transducers	Page 11
Dynamic Transducers	Page 13
Piezoelectric Transducers	Page 18
Magnetostrictive Transducers	Page 19
Impedance Matching of Transducers	Page 21
Appendix A—Logarithms	Page 25
Appendix B—Decibel Calculations	Page 28

*The more you know, the more you know you
ought to know.*

—Selected

TRANSDUCERS

Many electronic systems transmit, receive, count, record, or measure with the aid of **TRANSDUCERS**. Transducers take many forms and have a wide variety of applications, but basically they all convert energy from one form to another.

Transducers may be divided into two classes—input transducers and output transducers. **INPUT TRANSDUCERS** for electrically operated devices convert nonelectric energy such as sound, pressure or temperature to electric energy. This electric energy may be in the form of voltage or current. **OUTPUT TRANSDUCERS** for electrically operated devices work in the reverse order of input transducers in that they convert electric energy to forms of nonelectric energy. In some cases a single transducer may serve as both an input and an output device.

ENERGY CONVERTERS

Energy conversion is the process of converting one form of energy to another. There are many kinds of energy converters. However, some energy converters are called transducers, and some are not. For example, an electric toaster converts electricity to heat, but normally we do not think of a toaster as a transducer. Examples of common items that are called transducers are microphones and electric motors.

Transducers are useful because they can convert mechanical energy into electric energy, and then turn the electric energy back again into mechanical energy. To illustrate the use of transducers as energy converters, consider the block diagram in Figure 1. Starting at the left, energy of some form is fed to the input transducer. This energy may be mechanical pressure or motion, temperature, light or some other form of energy. The input transducer converts this nonelectric energy to electric energy.

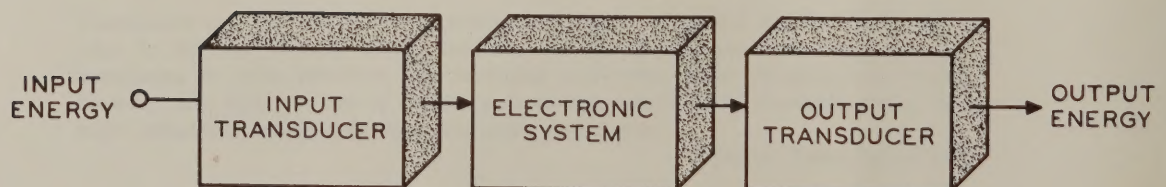


Figure 1

The electric energy developed by the input transducer becomes the input energy to the electronic system in Figure 1. The purpose of the electronic system is to amplify the useful electric energy to a desired magnitude, convert it to a desired frequency or reshape it to a desired waveform, and then route it to the output transducer.

The output of the electronic system is a form of electric energy, as is the input. Because the input and output are both electrical, there is no energy conversion in the electronic system. In this lesson we assume that the only device which changes energy from one form to another is a transducer.

The electrical output of the electronic system serves as the input energy to the output transducer. The output transducer converts this electric energy back into a form of nonelectric energy.

Figure 2 is a familiar application of the energy conversion principle—a public address (PA) system. Let's see how the various parts of this system fit into the block diagram of Figure 1. In the PA system, the input energy

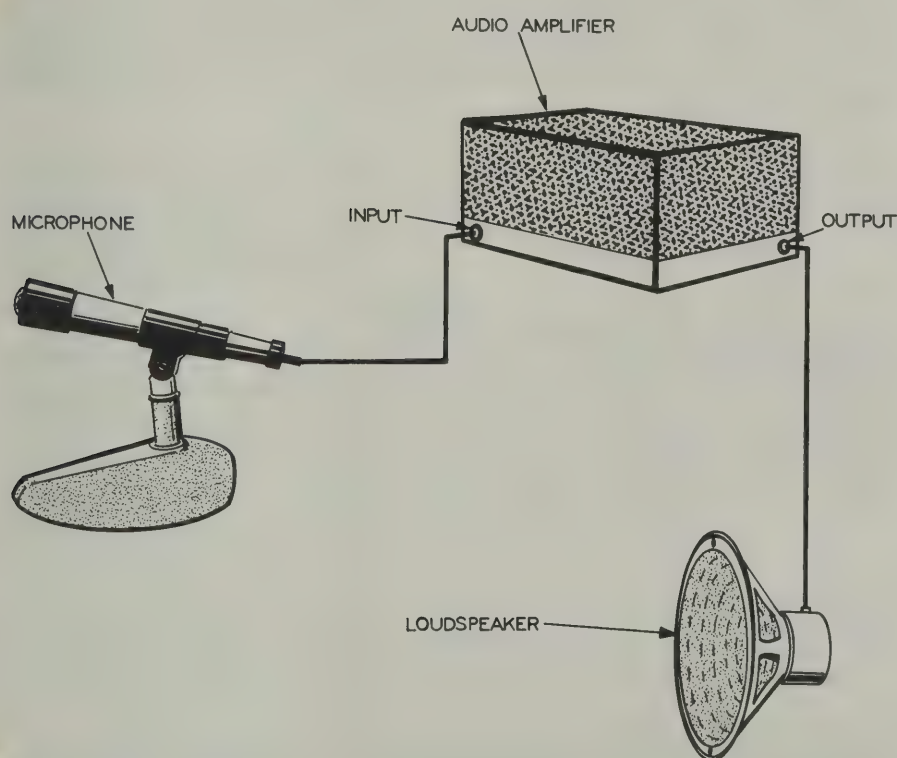


Figure 2

is the mechanical motion of sound waves (pressure)^h against the **MICROPHONE** head. A transducer in the microphone converts this mechanical motion to an electric signal. The electric signal is fed to an audio amplifier, which boosts the signal to a level that will operate the output transducer, a loudspeaker. The loudspeaker converts the electric signal to mechanical vibrations which produce the sound waves that reach our ears.

Input and output transducers are made in a variety of shapes and sizes and can be classified in several ways. One way to classify transducers is by their principles of operation. Based on this classification, the types of transducers to be considered in this lesson are variable resistance, variable reluctance, variable capacitance, dynamic, piezoelectric, and magnetostrictive.

Variable Resistance Transducers

VARIABLE RESISTANCE TRANSDUCERS are devices in which the energy impressed on the input of the transducer causes a change in its electrical resistance. For example, a potentiometer or rheostat can be used in determining fluid pressures in the arrangement of Figure 3. The electric circuit consists of rheostat R_1 in series with an electric current meter across a constant voltage E_s . R_1 serves as the input transducer which converts mechanical pressure displacements sensed by an air bladder. When air pressure is applied to the bladder, it inflates. This moves the wiper arm on the potentiometer, changing the circuit resistance and hence the output voltage of the circuit.

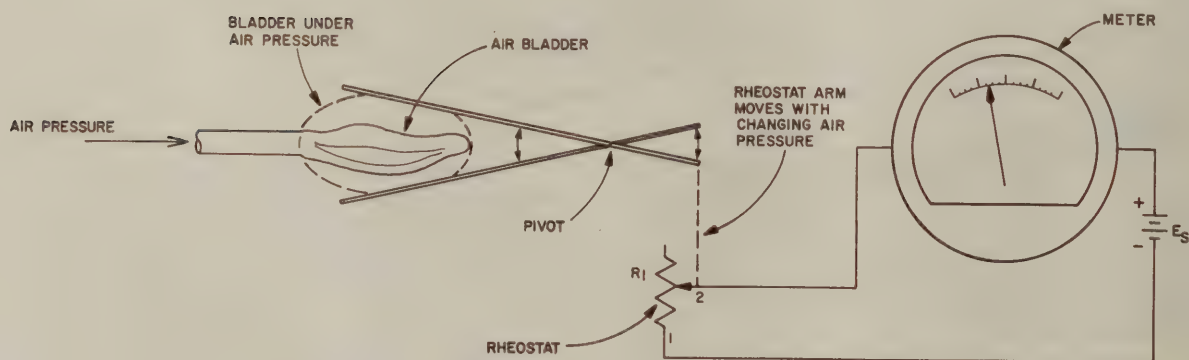
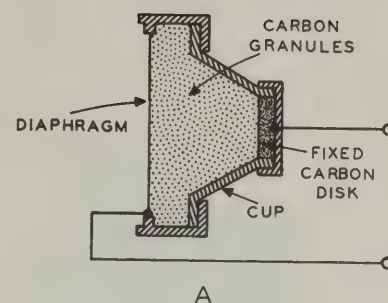


Figure 3

In Figure 3, R_1 itself does not provide an output voltage, but only varies the circuit resistance. It is therefore necessary to provide external voltage E_s . The electric meter serves as an output transducer which converts

electric current to a mechanical displacement of the meter pointer. The scale of the meter is calibrated in pressure units, such as pounds per square inch (PSI).

As air pressure in Figure 3 is increased, the arm of R_1 is forced downward, toward terminal 1 of R_1 . This results in a decrease of resistance between terminals 1 and 2 of R_1 which is in series with the meter. Ohm's Law states $I = E/R$. With a constant voltage E_s and a decrease in resistance, the current increases. In turn, the movement of the meter pointer up scale is proportional to the increase of current. This corresponds to the amount of pressure being measured.



A variable resistance transducer of another type is used in telephone microphone (mouthpiece) elements. This device is the **CARBON MICROPHONE** shown in Figure 4A. Basically, it consists of a cup-shaped housing with a fixed carbon disk attached to the narrow end of the cup. The cup is filled with carbon granules. The granules are sealed in by a flexible diaphragm at the wide end of the cup.

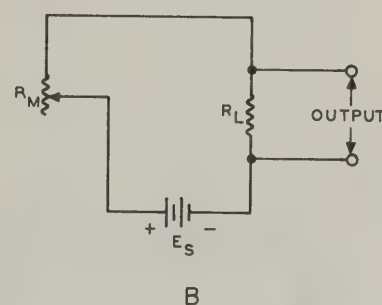


Figure 4

With normal air pressure acting on the diaphragm, the resistance of the carbon microphone remains constant. However, applying increased mechanical energy or pressure to the diaphragm causes the carbon granules to be packed closer together. This increased contact area between the diaphragm, carbon granules and fixed carbon disk results in a decrease of electrical resistance of the device.

When the added pressure on the diaphragm is removed, the diaphragm moves back to its original position. The contact area between the diaphragm, carbon granules, and fixed carbon disk decreases, causing the electrical resistance of the microphone to increase.

When sound waves strike the diaphragm of the microphone, the diaphragm vibrates at the frequency and magnitude of the sound waves. This causes corresponding changes in the resistance of the microphone. Thus, the sound waves applied to the diaphragm are converted to changes of electrical resistance. This action of the microphone can be compared to that of a variable resistor.

A carbon microphone can be connected into an electric circuit such as that of Figure 4B. Here the microphone is represented by the variable resistance R_M in series with load resistor R_L across a dc voltage source E_s .

As sound waves strike the diaphragm of the microphone, the resistance of the microphone varies according to the frequency and magnitude of the

TRANSDUCERS

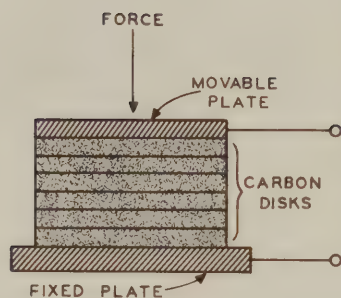
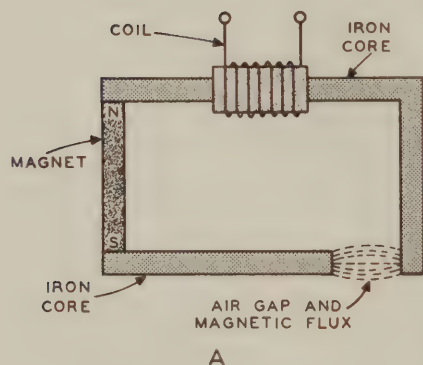


Figure 5

mechanical sound waves. The dc source, E_s , supplies electric energy for the microphone circuit. The changes of microphone resistance cause corresponding changes of circuit current which produce variations of voltage across R_L . This varying voltage across R_L is used as the electric output of the transducer.

The CARBON PILE shown in Figure 5 works on the same principle as the carbon microphone. The carbon pile has many applications, such as to control the speed of a motor, to measure shock and vibration of a structural member, or to measure the acceleration of a missile.

As a force or pressure is applied to the movable plate of the carbon pile, the contact area between the movable plate, carbon disks and fixed plate increases. This results in a decrease of electrical resistance between the plates of the carbon pile. As the amount of force or pressure is decreased, the contact area between the movable plate, carbon disks and fixed plate decreases. This causes the electrical resistance of the carbon pile to increase. The pile can be represented by a variable resistance such as that of Figure 4B.



Variable Reluctance Transducers

Mechanical energy can be changed to electric energy by a **VARIABLE RELUCTANCE TRANSDUCER**. In this type of transducer, mechanical energy causes a change in the reluctance of a magnetic circuit. The resultant change in the magnetic field causes an induced voltage which serves as the electric output of the transducer.

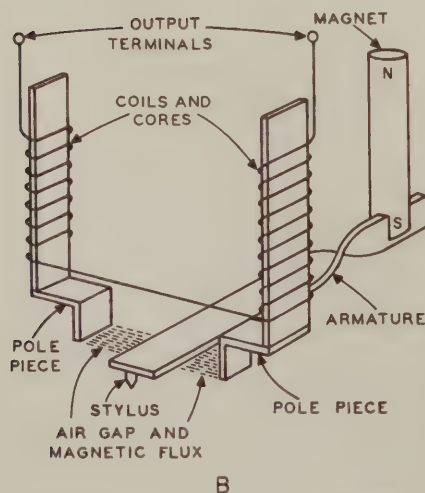


Figure 6

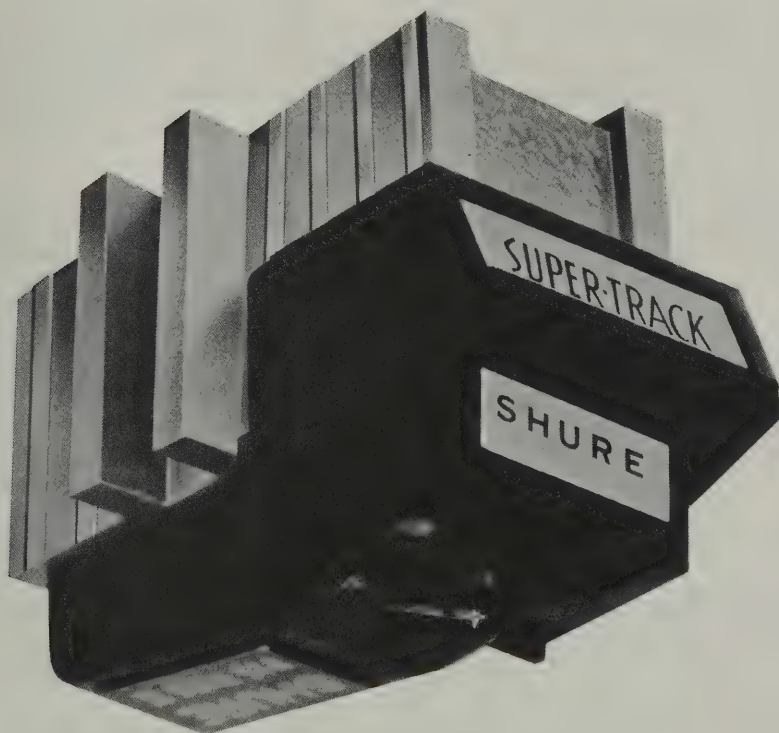
To better understand this principle, consider the magnetic circuit of Figure 6A. As shown, magnetic lines of force are concentrated in the magnet, iron cores and air gap. Changing the length of the air gap causes a change in the reluctance of the magnetic circuit. This increases or decreases the number of magnetic lines cutting the coil, and a small voltage is induced or generated in the coil.

One common application of this principle is in variable reluctance phonograph pickups. The basic elements that make up a variable reluctance cartridge used in a phonograph are shown in Figure 6B. One end of an armature is attached to the lower end of a magnet, while the free end of the armature, with the stylus, is able to move sideways toward the pole piece of either one of two coils. The side-to-side movement of the stylus and armature is caused by the variations in the walls of the record groove.

The magnetic path in Figure 6B is from the north pole of the magnet through the air to the cores of both coils, through the cores and pole pieces to the armature, and back along the armature to the south pole of the magnet. As long as the stylus is motionless, there is no change in the magnetic circuit, the field is steady, and there is no voltage induced in either coil. As the stylus moves sideways in the record groove, it causes the armature to move in the same direction. Thus, the air gap between the armature and one pole piece is decreased, while the air gap between the armature and the other pole piece is increased.

For example, suppose the groove variations cause the stylus and armature in Figure 6B to move toward the left pole piece. This shortens the air gap between the armature and left pole piece and reduces the reluctance in this portion of the magnetic circuit. At the same time, the gap width increases between the armature and right pole piece, increasing the reluctance in this part of the magnetic circuit. Since magnetic flux varies inversely as reluctance, the number of magnetic lines of force threading through the left coil increases, while the number of magnetic lines through the right coil decreases.

A voltage is induced whenever a coil is cut by a changing magnetic field. In this example, the expanding field around the left core induces a voltage in



*A high-quality magnetic phonograph cartridge capable of tracing complex record groove modulations with less than one gram of tracking pressure.
Courtesy Shure Bros., Inc.*

the left coil. At the same time, the field around the right core is collapsing due to the increased reluctance in its path. Thus, a voltage is induced in the right coil, but of a polarity opposite to that in the left coil. The coils are connected in series and are wound in such a way that the voltage induced in one aids the voltage induced in the other. This is called a series aiding connection. With this arrangement, the total output of the cartridge is approximately twice that of a single coil.

Because of its mechanical design, vertical movement of the armature of Figure 6B is limited within the area of the pole pieces. Therefore, a vertical movement of the armature has no appreciable effect on the reluctance and there is no changing magnetic field to induce voltage in the coils.

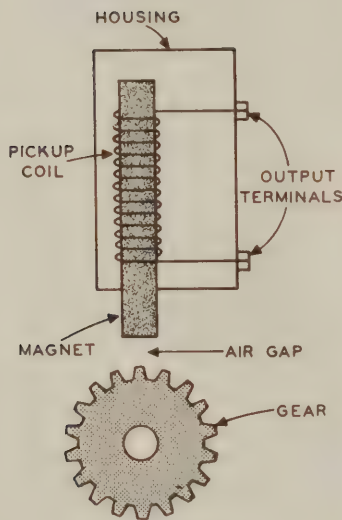


Figure
7

The magnetic pulse detector shown in Figure 7 is another form of variable reluctance transducer. This type of transducer is useful in counting shaft revolutions in machinery, in measuring propeller speeds of ships and in measuring fluid flow for industry.

An example of this type of transducer is a pickup coil wound on a bar magnet and placed near a revolving toothed gear. The magnetic lines of force are concentrated in the bar magnet, air gap and toothed gear. When a tooth passes the pickup coil and magnet, the air gap between them decreases. This reduces the reluctance and strengthens the magnetic field. As a result, more magnetic lines of force cut the coil turns, generating a voltage pulse in the coil.

This voltage pulse may be transferred to an amplifier. Often, the output pulse from this type of detector is strong enough to drive an electronic counter directly.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

ENERGY CONVERTERS

VARIABLE RESISTANCE TRANSDUCERS

VARIABLE RELUCTANCE TRANSDUCERS

1. How does an input transducer differ from an output transducer?
2. Microphones and electric motors may be called transducers. One is an input transducer and one is an output transducer. Which one is the input transducer?
3. The electronic system of Figure 1 can serve only to amplify, convert to another frequency, or reshape the waveform of electric energy. True or False?
4. The microphone in the PA system of Figure 2 converts mechanical vibrations of sound waves to proportional changes of electric energy and therefore can be considered as (a) an output transducer, (b) an input transducer.
5. The audio amplifier of Figure 2 is an electronic system which amplifies the signal from the microphone to a level sufficient to operate the loudspeaker. The input and output energy of the audio amplifier are (a) both electric, (b) both nonelectric.
6. The loudspeaker of Figure 2 converts the amplified audio signal from the amplifier to sound energy and thus can be considered (a) an output transducer, (b) an input transducer.
7. A variable resistance transducer is a device in which the input energy changes cause corresponding changes in the electrical (a) capacitance, (b) inductance, (c) resistance.
8. The variable resistor of Figure 3 serves as (a) an input transducer, (b) an output transducer.
9. What electric device serves as an output transducer in Figure 3?
10. Increasing the contact area between the diaphragm, carbon granules, and fixed carbon disk of a carbon microphone will decrease the electrical resistance of the microphone. True or False?
11. The electric action of a carbon microphone can be compared to that of (a) a capacitor, (b) an inductor, (c) a variable resistor.

12. The carbon microphone generates its own voltage and does not require an external source of energy. True or False?
13. The carbon pile works on the same principle as the _____.
14. In Figure 6A, changing the length of the air gap results in a corresponding change of the magnetic circuit's (a) capacitance, (b) reluctance.
15. Explain how a voltage is generated in the coil of the transducer of Figure 6A.
16. What causes the side-to-side movement of the stylus and armature of a variable reluctance cartridge?
17. The stylus of the variable reluctance cartridge of Figure 6B is attached to the armature's (a) fixed end, (b) free end.
18. In Figure 6B, as the armature moves toward the left pole piece, the reluctance between the armature and right pole piece (a) increases, (b) decreases.
19. The two coils of the variable reluctance cartridge in Figure 6B are series (a) opposing, (b) aiding.

Variable Capacitance Transducers

VARIABLE CAPACITANCE TRANSDUCERS work on the principle that energy impressed on the input causes a corresponding change in capacitance. This idea is used in the **CAPACITOR MICROPHONE** shown in Figure 8.

A capacitor is formed by the metallic diaphragm and the fixed metal plate of the microphone. These two pieces act as plates of the capacitor. Air forms the dielectric between the two plates.

As sound waves strike the diaphragm, the diaphragm vibrates according to the frequency and magnitude of the sound waves. This vibration of the diaphragm causes corresponding **CHANGES IN DISTANCE** between the plates, and thus **VARIES THE CAPACITANCE** of the microphone.

The capacitor microphone does not generate a voltage itself. It only varies the capacitance of the circuit to which it is connected. Therefore it is necessary to connect the microphone as shown in Figure 8. Here it is in series with a constant high voltage E_s and fixed load resistor R_L . Voltage E_s is impressed across the diaphragm and fixed plate of the microphone through resistor R_L . This causes the microphone to take a charge proportional to its capacitance. As the diaphragm moves in and out, the capacitance of the microphone changes. This causes a change in the charging current through fixed resistor R_L and the voltage across R_L . This voltage across R_L is the output of the microphone. The output voltage is fed into an amplifier where it is amplified to sufficient magnitude to drive or operate a load.

The can alignment detector shown in Figure 9 works on the principle of varying capacitance. A capacitive probe forms the fixed plate of the capacitor and is connected directly to the input of the control system. The cans moving along the conveyor form the movable plate and are connected to the control system through a ground connection. Air serves as a dielectric between the plates.

An improperly positioned can or a can of the wrong size moving along the conveyor will change the capacitance of the system. If a can is in the wrong position on the belt, the distance between the plates is varied, causing a change in circuit capacitance. Also, if the can in front of the probe is larger or smaller than the rest of the cans moving along the conveyor, the circuit capacitance is varied.

With a properly aligned can of the right size moving along the conveyor, the control system is set at a desired capacitance to maintain a constant speed of the conveyor.

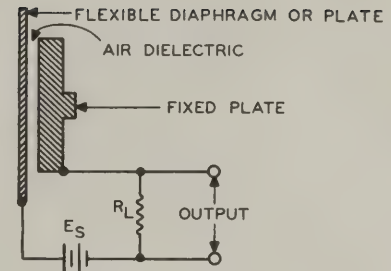


Figure 8

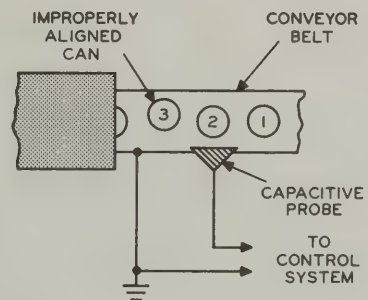


Figure 9



Capacitor microphones are well suited to professional use in radio, TV, and recording studios.

Courtesy North American Philips Co., Inc.

In Figure 9, the conveyor moves can number 1 in front of the probe. The probe senses the can as being of correct size and distance from the probe. Thus, the proper capacitance of the control system is maintained, the conveyor keeps moving, and can 1 passes the probe.

As can 1 passes the probe, can 2 is moving into position in front of the probe. The probe now senses can 2 as being of correct size and position, so the conveyor moves the next can in front of the probe.

Can 3 is not positioned correctly. As can 3 moves in front of the probe, the probe detects a change of distance between it and can 3. This distance changes the capacitance of the control system. This change of capacitance sends a voltage pulse to the control switch to stop the conveyor until the can is realigned, or rejected.

If can 3 is either larger or smaller than cans 1 and 2, it represents a change in plate area of the system, causing a corresponding change in circuit capacitance. Once again, the change of capacitance sends a voltage pulse to the control switch to stop the conveyor until the can is rejected or replaced.

This principle of the variable capacitance transducer is accomplished either by changing the distance between the plates, plate area, or dielectric of the device.

These principles are also applied in determining the level of liquid in a tank, operating a warning signal for overflow in a container, or for turning off power in machinery to prevent an accident to the operator.

Dynamic Transducers

When a mechanical force or pressure causes a coil to vibrate or move in a magnetic field, a voltage is induced in the coil. Thus, mechanical energy is converted to electric energy. In like manner, when a varying voltage is applied to the terminals of the same coil, the coil will vibrate, or move mechanically. This vibration depends on the frequency and magnitude of the applied voltage. Thus, electric energy is converted to mechanical energy.

The two physical effects just described are the basis of **DYNAMIC TRANSDUCERS**. This type of transducer can serve as either an input or an output device.

Most **LOUDSPEAKERS** commonly found in radios, television receivers, and sound systems are dynamic output transducers. These loudspeakers are classified as either permanent magnet (PM) or electromagnetic (EM) types.

The basic construction of a PM dynamic loudspeaker is shown in Figure 10A. The **POLE PIECE**, which consists of a small cylindrical magnet with a soft iron wafer around it, is friction fitted into a large rectangular piece of soft iron labeled the **POT**. This forms the permanent magnet of the speaker and provides a complete magnetic path for the pole piece.

The pole piece extends partly through the junction of the pot and metal frame. Around the pole piece, but not touching it or the pot, is a hollow form with a coil of wire on it. This coil is the **VOICE COIL**, which receives the electric energy from the output of the electronic system. The voice coil is centered around the pole piece by a flexible fiber or metal

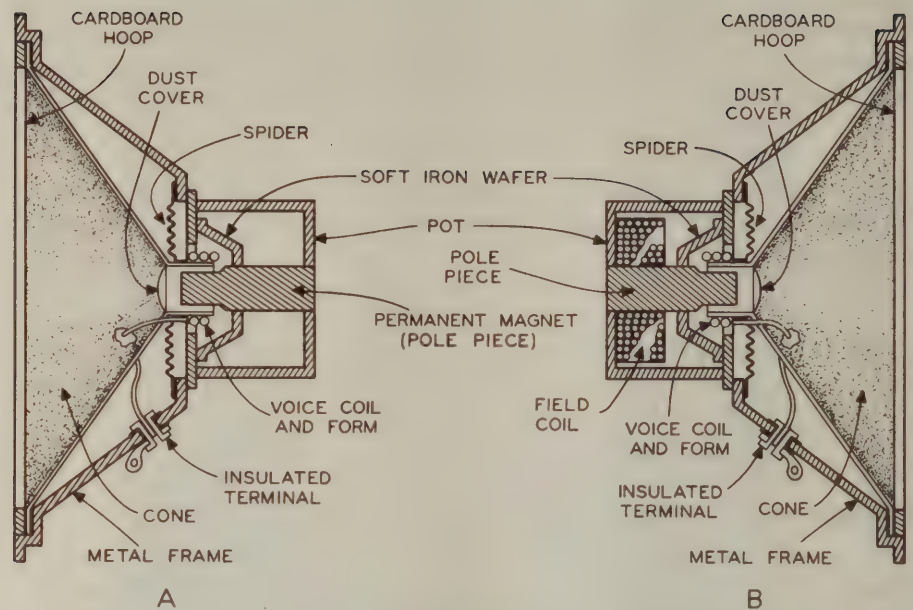


Figure 10

washer or elasticized cloth called a **SPIDER**. The spider holds the voice coil in its normal position when there is no signal current in the voice coil, and allows the voice coil to move back and forth along the pole piece when a signal is applied.

The voice coil and attached spider are cemented to the **SPEAKER CONE**. The speaker cone is a cone-shaped diaphragm made of pressed paper. The outer rim of the cone is cemented to the metal frame and held tight all around by a cardboard or cork hoop or ring. Motion of the voice coil causes the cone to vibrate and create sound waves in the air.

The fragile voice coil wires are brought out and cemented to the surface of the speaker cone. Strong flexible leads are soldered to the voice coil wires at this point. The other ends of the flexible leads are connected to insulated terminals mounted on the frame of the loudspeaker.

A **DUST COVER**, located at the junction of the speaker cone and voice coil form, covers the air gap between the voice coil and pole piece. This prevents dirt, dust, pieces of metal or any other foreign particles from entering the air gap and affecting the mechanical action of the speaker.

The construction of an electromagnetic (EM) loudspeaker is shown in Figure 10B. It differs from the PM speaker mainly in that the pole piece is not a permanent magnet. The pole piece of the EM speaker is made of soft iron and has a large coil of wire called a FIELD COIL around it. This arrangement serves as an electromagnet.

The field coil is energized by a dc current supplied by the power supply of the electronic system. When the field coil becomes energized, a steady magnetic field is produced around the coil and pole piece. The pole piece then becomes magnetized and serves the same purpose as the permanent magnet of the PM speaker.

When a signal current passes through the voice coil in one direction, it produces a magnetic field which opposes the fixed magnetic field of the PM or EM speaker. When the current is reversed, the direction of the voice



A siren converts electric energy to a shrill audible tone which can be heard over great distances. A motor is used to turn a perforated wheel or drum at high speed which produces powerful sound waves.

Courtesy Federal Sign and Signal Corp.

TRANSDUCERS

coil field is reversed so that it aids the steady field of the magnet. Thus, current in the voice coil causes the coil to be either attracted or repelled by the magnet.

Since the cone is attached to the coil, vibrations of the voice coil cause corresponding vibrations of the cone.

The dynamic microphone of Figure 11 works on the same principle but in the reverse order of the loudspeaker. As sound waves strike the diaphragm, the diaphragm and attached coil move back and forth within the magnetic field. As the coil moves in the magnetic field, a voltage is induced in the coil. This voltage depends on the frequency and magnitude of the sound waves which move the coil.

Since the dynamic microphone actually generates a voltage, no external voltage source is necessary. However, the output signal of the microphone is not sufficient to operate the output transducer directly. Therefore, it is necessary to amplify the output signal by means of an amplifier.

Since the operation of a dynamic microphone is the reverse of that of a dynamic loudspeaker, it is possible to use a speaker as a microphone. This is often done in intercommunication systems commonly found in factories, homes, and aboard ships.

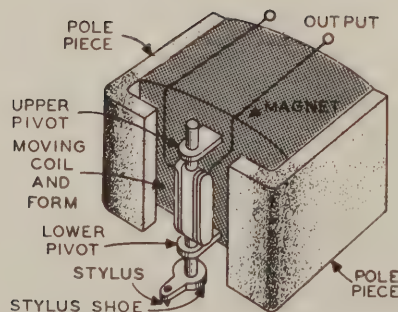


Figure 12

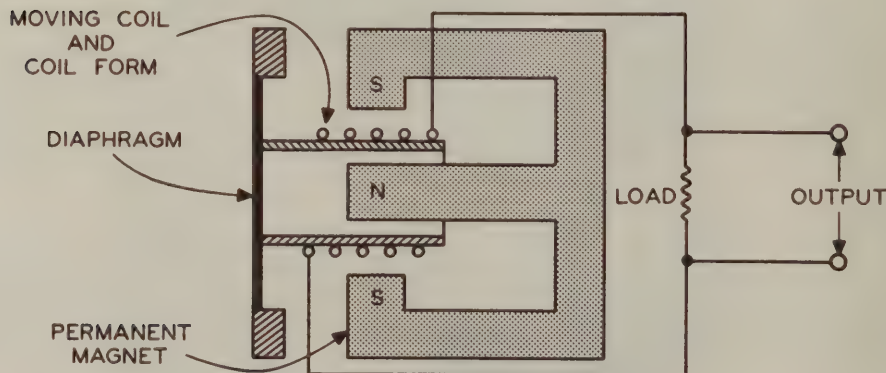


Figure 11

Another dynamic transducer is the moving coil phonograph cartridge shown in Figure 12. The moving coil is wound on a long, narrow form which

has pivots located at its upper and lower ends. The coil is placed between the pole pieces of a strong magnet. Attached to the lower end of the moving coil form is a stylus shoe which in turn holds the stylus.

The side-to-side motion of the stylus, produced by the variations in the walls of the record groove, causes the small coil of wire to rotate on its vertical axis. The coil cuts the lines of force created by the magnetic field. As a result, a small signal voltage is induced in the coil. This voltage is then fed from the output terminals to an amplifier.

Piezoelectric Transducers

When certain types of crystals such as quartz, tourmaline, or Rochelle salt are twisted, bent or compressed by a mechanical force, a voltage is generated. This action is known as the **PIEZOELECTRIC EFFECT**, and the electricity generated is known as piezoelectricity.

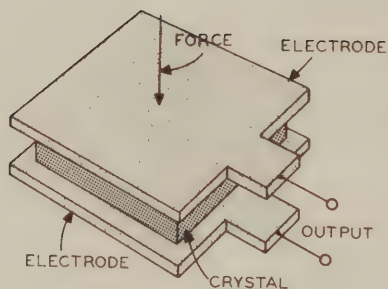


Figure 13

A piezoelectric or crystal transducer is shown in Figure 13. A voltage is developed between the opposite faces of the crystal when it is stressed mechanically. Therefore, the device can act as a voltage generator without any externally applied voltage. On the other hand, when a varying voltage is impressed across the electrodes attached to the opposite faces of the crystal, the crystal is disturbed and it produces mechanical energy or motion. Thus, some crystal transducers are input transducers, and others are output transducers.

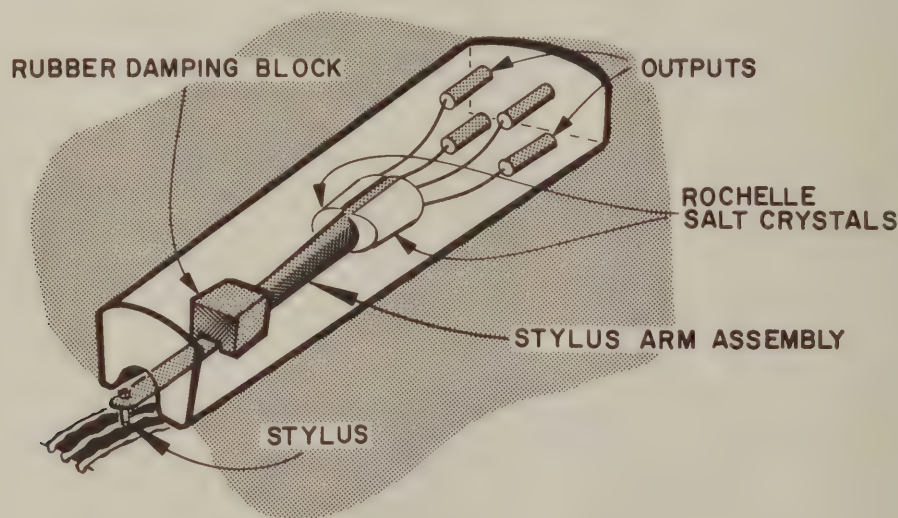


Figure 14

A practical application of a piezoelectric transducer is the crystal phonograph cartridge of Figure 14. The mechanical vibrations caused by the grooves of a record mechanically vibrate the stylus of the crystal cartridge. The damping block, when used, prevents excessive stylus travel. These vibrations of the stylus exert a varying pressure on the crystal. This causes a varying output voltage, which is then fed into an amplifier.

Another application of the piezoelectric effect is the measurement of force and acceleration by the use of an **ACCELEROMETER** such as shown in

The following Practice Exercise questions cover the subjects which you have just studied. They are:

ENERGY CONVERTERS

VARIABLE CAPACITANCE TRANSDUCERS

DYNAMIC TRANSDUCERS

20. Both the carbon microphone and the capacitor microphone require an external voltage source. True or False?
21. Explain how the output voltage of the capacitor microphone is developed.
22. Dynamic transducers serve only as output devices. True or False?
23. What are two classifications of dynamic loudspeakers?
24. A loudspeaker used in a radio receiver serves as (a) an output transducer, (b) an input transducer.
25. To produce its steady magnetic field, an EM speaker uses (a) a permanent magnet, (b) an electromagnet.
26. The voice coil of a loudspeaker vibrates when (a) alternating current passes through it, (b) direct current passes through it.
27. A speaker using a voice coil moving in a magnetic field is called a _____ speaker.
28. A dynamic microphone used in a public address system works in the reverse order of the dynamic loudspeaker used in the same system. True or False?
29. A dynamic microphone does not generate a voltage directly, and it is therefore necessary to use an external voltage source. True or False?

Figure 15. This transducer works on the same principle as the one shown in Figure 13. A downward force on the face plate causes the crystal to be stressed mechanically. When the crystal is in this stressed condition, a voltage is generated across the opposite faces of the crystal, between the output terminals. As the downward force increases, the pressure on the crystal increases, causing an increase in the output voltage.

Decreasing the downward force decreases the pressure on the crystal and therefore reduces the output voltage. When the downward force is decreased or eliminated, the spring moves the face plate back to its normal position.

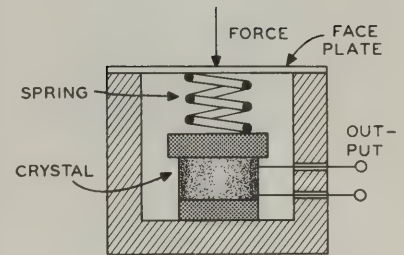


Figure 15

A limitation of this transducer is that it responds only to changes of pressure. The voltage generated by an applied steady pressure decays or leaks off. Unless the crystal is connected to an amplifier of infinite impedance, the voltage produced by the permanent force or pressure on the crystal soon leaks away. As a result, there is no steady voltage to indicate the permanent force or displacement.

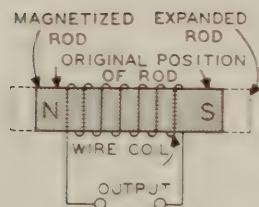
One application of this device is its use in connection with a shock table for shock and vibration testing of various types of equipment. For example, electronic equipment used in commercial aircraft or missiles must withstand a certain amount of shock and vibration under normal operating conditions. As part of the final tests, the equipment undergoes shock and vibration tests at the factory where it is manufactured. The accelerometer is attached to the shock machine on which the equipment is tested. As the shock machine vibrates, the accelerometer measures and records the amount of shock and vibration the equipment can withstand, and at what point the equipment fails.

A crystal will vibrate more freely at some given frequency than at other frequencies, depending upon its size and shape. This is known as the resonant frequency of the crystal. Thus, an applied voltage at the resonant frequency will cause a greater amount of vibration than will an equal amount of voltage at a frequency above or below resonance. Also, mechanical vibrations at the resonant frequency of a crystal will produce a greater output voltage than equal amplitude vibrations at some frequency above or below the mechanical resonant frequency. This property is useful because it permits crystals to be used as frequency selective and oscillator elements in some control systems.

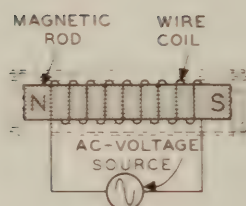
Magnetostrictive Transducers

A **MAGNETOSTRICTIVE TRANSDUCER** is a device in which a magnetized rod placed inside a coil of wire induces a voltage in the coil when the

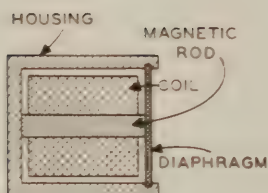
TRANSDUCERS



A



B



C

Figure 16

rod is elongated or compressed. For example, when the rod of Figure 16A is elongated or lengthened, the magnetic lines of force of the magnet cut across the turns of the coil, inducing a voltage of one polarity in the coil. When the rod returns to its original position, or when it is compressed, the magnetic lines of force induce a voltage of opposite polarity in the coil. Thus, an ac voltage is developed across the terminals of the coil. This device converts mechanical energy to magnetic energy, and magnetic energy to electric energy. Used in this arrangement, the device is an input transducer.

Transducers based on the principle of Figure 16A are used to determine the amount of vibration in structural members of aircraft and missiles. As the structure vibrates, the vibrations cause the magnetic rod of the transducer to vibrate at the same frequency. In turn, the magnetic lines of force cut across the turns of the coil, inducing an ac voltage in the coil. The output voltage is fed to a meter or recorder which tells the technician how much vibration the structure is under.

A magnetostrictive transducer is arranged to serve as an output transducer in Figure 16B. With an ac voltage impressed across the terminals of the coil, the coil sets up its own magnetic field. This varying magnetic field causes the rod to change its length or vibrate at a frequency and magnitude proportional to the applied ac voltage.

A more practical arrangement of the magnetostrictive transducer is shown in Figure 16C. The rod and coil are placed in a cup-shaped housing. One end of the rod is fastened rigidly to the closed end of the housing. A flexible diaphragm is attached to the open end of the housing and free end of the rod.



This depth sounder converts electric energy to ultrasonic sound waves, detects the reflected (bounced) wave, and converts the sound back into an electric energy readout. The transducer head mounts below the water line while the main unit provides the readout in the cabin.

Courtesy Heath Company

A high frequency ac voltage applied to the coil causes the rod and diaphragm to vibrate. A typical application of this arrangement might be that of mixing paint. For example, by placing the face of the diaphragm in contact with the paint in a large container, the paint is agitated or mixed when a high frequency ac voltage is applied to the transducer.

The widest applications of magnetostrictive transducers are in the field of **ULTRASONICS**. Ultrasonics deals with the production, reception, detection, and absorption of acoustic waves at frequencies of 20 kHz or above. This field is finding widespread use in modern industry.

IMPEDANCE MATCHING OF TRANSDUCERS

A transducer operates best when its impedance matches that of the system to which it is connected. Mismatching the input or the output impedance of a device will usually introduce distortion of the signal or result in a drop in signal level. **MAXIMUM TRANSFER OF POWER** occurs when the load impedance matches the source impedance.

Careful consideration must also be given to matching the input and output voltages of the various components and systems. If the output voltage of a device is too high for the input of the system to which it is applied, serious overloading and distortion may result in the system. Often the equipment itself may be damaged. On the other hand, if too little voltage is applied, it may not be possible to obtain a suitable output from the system.

For example, the output energy of a dynamic microphone is not sufficient to operate a speaker directly. It is therefore frequently necessary to connect the microphone through a suitable microphone transformer to an amplifier. The amplifier amplifies the energy to a level that will operate a speaker.

The internal impedance of a typical dynamic microphone is about 150 ohms, whereas an amplifier may have an input impedance of about 200,000 ohms. To obtain maximum power transfer it is necessary to match the microphone impedance with its load, which in this case is the input impedance of the amplifier. This is done by connecting an impedance matching transformer called a **MICROPHONE TRANSFORMER** between the microphone and input of the amplifier, as shown in Figure 17. Most microphones contain an impedance matching transformer.

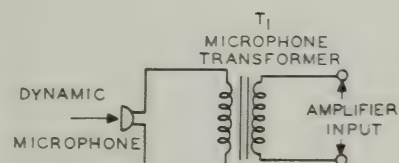


Figure 17

The primary impedance of transformer T_1 in Figure 17 matches the impedance of the microphone. The secondary impedance of T_1 matches the input

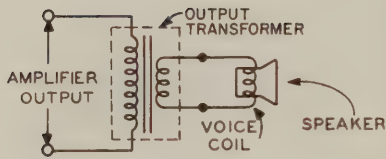


Figure
18

impedance of the amplifier. It can be said that the microphone “sees” the primary impedance of T_1 as a matching impedance. The secondary, which by transformer action receives the power from the primary, “sees” the input impedance of the amplifier as a matching impedance. Therefore, the microphone transformer matches the low impedance source to the high impedance load.

The problem of matching the OUTPUT of an electronic system to the output transducer is also solved by a matching transformer. For example, in a radio receiver the output power is fed through an **OUTPUT TRANSFORMER** to the voice coil of the loudspeaker, as shown in Figure 18. The impedance of the voice coil is generally in the order of 3 to 16 ohms. An amplifier using vacuum tubes must have a load of thousands of ohms. Therefore, the output transformer is used to couple the output of the amplifier to the voice coil of the loudspeaker to obtain maximum transfer of energy.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

ENERGY CONVERTERS

PIEZOELECTRIC TRANSDUCERS

MAGNETOSTRICTIVE TRANSDUCERS

IMPEDANCE MATCHING OF TRANSDUCERS

30. When certain types of crystals are bent or compressed by a mechanical force, a voltage is generated between the opposite faces of the crystal. This effect is commonly called the _____ effect.
31. When used as an input device, such as a crystal microphone, the crystal transducer acts as a voltage generator. True or False?
32. When a varying voltage is impressed across the opposite faces of a crystal transducer, what physical effect takes place?
33. The output of the magnetostrictive transducer of Figure 16A is (a) dc, (b) ac.
34. A piezoelectric crystal will vibrate more readily at a particular frequency than at some other frequency. This is the _____ frequency.
35. A magnetostrictive transducer is only capable of converting mechanical energy to electric energy. True or False?
36. The widest application of magnetostrictive transducers is in the field of communications. True or False?
37. Impedance matching is a condition in which the impedance of a component or device is equal to another impedance to which it is connected. True or False?
38. Maximum transfer of power occurs when the load impedance is equal to one-half that of the source impedance. True or False?
39. It can be said that a dynamic microphone sees the secondary impedance of a transformer as a matching impedance. True or False?

SUMMARY

Basically, a transducer converts energy from one form to another, and has an output proportional to its input. There are two groups of transducers; input transducers and output transducers. Input transducers change non-electric energy to electric energy. Output transducers change electric energy to nonelectric energy. In some cases a single transducer is capable of serving both as an input and output device.

Transducers are also classified according to their operating principles. For example, the variable resistance transducer is a device in which energy impressed on its input causes corresponding changes of its electrical resistance. It requires an external voltage source, since it does not generate its own voltage. The variable reluctance device is another form of input transducer. It converts mechanical energy to electric energy by changing the reluctance of a magnetic circuit. The variable reluctance transducer generates its own voltage and does not require an external voltage source.

The variable capacitance transducer is another device for converting energy. Nonelectric energy impressed on the input of this device will cause a corresponding change in its electrical capacitance. Like the variable resistance transducer, this device requires an external voltage source and amplification of its output energy.

The dynamic transducer can serve as either an input or an output device. In this device, a mechanical force exerted on the coil causes the coil to move in a fixed magnetic field. As the coil cuts the lines of force, a voltage is generated in the coil. Also, by passing an alternating current through the coil, a varying magnetic field is developed about the coil. Thus, the coil interacts with the fixed magnetic field, causing the coil to vibrate mechanically. In some cases, a dynamic transducer serves as a combination input-output transducer, such as a speaker-microphone of an intercom system. It generates its own voltage and therefore does not require an external voltage source.

Piezoelectric transducers are certain types of devices containing crystals that generate an electric voltage whenever they are bent or compressed. Also, when an ac signal is placed across the opposite faces of the crystal, the crystal vibrates in proportion to the frequency and magnitude of the ac signal. This principle is applied to both input and output transducers.

By elongating and compressing a magnetized rod in a coil of wire of a magnetostrictive transducer, an ac voltage is induced into the coil. Also, by placing a high frequency ac voltage across the terminals of the coil, the rod will vibrate mechanically in proportion to the frequency and amplitude of the

ac signal as shown by the dashed lines around the bar in Figure 16B. The magnetostrictive transducer serves as an input and an output device and is widely used in the field of ultrasonics.

Often the output signal of a transducer is not of the desired amplitude, frequency, or waveform to operate its load. Thus, the output of a transducer is often fed into an electronic system where the signal is changed to the desired magnitude, frequency, or waveform. Impedance matching devices are necessary to transfer the maximum power from the transducer to the electronic system or from the electronic system to the transducer. Without them, mismatching may occur, causing distortion or a drop in signal level. Input and output voltages should also be taken into consideration. If the output voltage is too high for the input of a system, serious overloading and distortion will result. In some cases the equipment itself may be damaged.

APPENDIX A — LOGARITHMS

To simplify the writing of arithmetic problems, it is customary to use exponents. These are small numbers written slightly above and to the right of a larger number. Thus, the number 4 with an exponent 2, or 4^2 , is read as "Four Squared" and indicates that 4 is to be used twice as a factor. As an equation:

$$4^2 = 4 \times 4 = 16$$

The same plan is true for all numbers and exponents and using the number 10, as further examples:

$$10^0 = 1$$

$$10^1 = 10$$

$$10^2 = 10 \times 10 = 100$$

$$10^3 = 10 \times 10 \times 10 = 1,000$$

$$10^4 = 10 \times 10 \times 10 \times 10 = 10,000$$

Exponents are read also as the power to which a number is raised. Thus:

10^2 indicates 10 raised to the second power.

10^3 indicates 10 raised to the third power.

10^4 indicates 10 raised to the fourth power.

Exponents 0 and 1 are seldom used because they are not necessary. 10^0 indicates no use as a factor and thus is equal to 1 or unity which does not affect any multiplication. 10^1 which indicates ten raised to the first power is not needed since it is used once as a factor without the exponent. However, 10 with exponents between 0 and 1 must equal numbers between 1 and 10. In the same way, 10 with exponents between 1 and 2 must equal numbers between 10 and 100 and so on.

As a definition, THE COMMON LOGARITHM OF A NUMBER IS THE POWER TO WHICH 10 MUST BE RAISED TO EQUAL THAT NUMBER. Accordingly, the previous table can be rewritten as:

$$\log 1 = 0$$

$$\log 10 = 1$$

$$\log 100 = 2$$

$$\log 1,000 = 3$$

$$\log 10,000 = 4$$

APPENDIX A (Continued)

To illustrate intermediate numbers, assume 10 with an exponent of $\frac{1}{2}$ or .5. This indicates that 10 is raised to the $\frac{1}{2}$ power which is a number that multiplied by itself will equal 10. In arithmetic this is known as the square root, which for 10 is 3.162. In the form of an equation:

$$10^{\frac{1}{2}} = 10^{.5} = 3.162, \text{ or } \log 3.162 = .5$$

Further calculations of this kind are not necessary because the charts at the end of this lesson list the logarithms of all numbers between 1 and 10 to an accuracy of two decimal places. Referring to these charts, the left hand "N" column lists numbers which reading downward increase in value by tenths. Starting at the left again, the remaining ten columns are headed by a number which indicates the hundredths of those in the N column.

To check the example problem:

$$\log 3.162 = .5$$

follow down the N column to 3.1. Then across to the column headed "6" which represents .06 because $3.1 + .06 = 3.16$. The 3.1 line in the 6 column lists the logarithm as .4997. Therefore, according to the chart:

$$\log 3.16 = .4997$$

Thus, the chart is accurate to two decimal places which is sufficient for most ordinary work. For practice, check the following logs on the chart.

$$\log 1.75 = .2430$$

$$\log 5.29 = .7235$$

$$\log 2.37 = .3747$$

$$\log 6.33 = .8014$$

$$\log 3.72 = .5705$$

$$\log 7.98 = .9020$$

$$\log 4.00 = .6021$$

$$\log 9.11 = .9595$$

Going back to the previous explanations if—

$$10^5 = 3.162 \text{ then:}$$

$$10^{1.5} = 10 \times 3.162 = 31.62$$

$$10^{2.5} = 10 \times 10 \times 3.162 = 316.2$$

$$10^{3.5} = 10 \times 10 \times 10 \times 3.162 = 3,162$$

$$10^{4.5} = 10 \times 10 \times 10 \times 10 \times 3.162 = 31,620$$

APPENDIX A (Continued)

With the exception of the first, each of these exponents consists of a whole number and the same decimal fraction. Note carefully, when the whole number of the exponent is increased by one, the resulting number is increased ten times. Thinking of these exponents as logarithms, the whole number is the CHARACTERISTIC and the decimal fraction is the MANTISSA. Thus, the charts list only the Mantissas because the logarithms extend from 0 to .9996 for numbers between 1 and 9.99. However, the decimal relationship makes the charts equally useful for higher numbers.

By moving the decimal point and indicating the corresponding powers of 10, any number can be written in reduced form. Reversing a previous list:

$$3.162 = 3.162 \times 10^0 = 10^{-.5}$$

$$31.62 = 3.162 \times 10^1 = 10^{1.5}$$

$$316.2 = 3.162 \times 10^2 = 10^{2.5}$$

$$3,162. = 3.162 \times 10^3 = 10^{3.5}$$

$$31,620. = 3.162 \times 10^4 = 10^{4.5}$$

In effect, each time the decimal point in the left column is moved one place to the right, the characteristic of the exponent in the right column is increased by "1". In all cases, the mantissa remains the same. Thus, the characteristic of a logarithm is equal to the number of places the decimal point is moved to the left to reduce a number to less than "10". The mantissa of the logarithm of the reduced number is listed in the charts.

For example, to find the logarithm of 452, the decimal point is moved "2" places to the left to reduce the number to 4.52. According to the chart, the mantissa of 4.52 is .6551. Therefore:

$$\log 452 = 2.6551.$$

In like manner:

$$\log 37 = 1.5682$$

$$\log 506 = 2.7042$$

$$\log 1500 = 3.1761$$

APPENDIX B — DECIBEL CALCULATIONS

The decibel equation:

$$\text{dB} = 10 \log \frac{P_2}{P_1}$$

indicates a gain of power as long as the output P_2 is greater than the input P_1 .

In some cases, there is a loss or attenuation, and to simplify the calculations, the power ratio is inverted and a minus sign is added. As a loss equation:

$$\text{dB} = -10 \log \frac{P_1}{P_2}$$

By this method, all ratios are greater than "1" and can be calculated in exactly the same way.

As an example, suppose an amplifier increases the power from 1 watt to 50 watts. In decibels, the gain is:

$$\begin{aligned} \text{dB} &= 10 \log \frac{P_2}{P_1} = 10 \log \frac{50}{1} \\ &= 10 (\log 50) = 10 \times 1.6990 \\ &= 16.99 \text{ or approximately } 17. \end{aligned}$$

If, for some reason, the power is reduced from 50 watts back to 1 watt, the loss is:

$$\begin{aligned} \text{dB} &= -10 \log \frac{P_1}{P_2} \\ &= -10 \log \frac{50}{1} \\ &= -10 (\log 50) \\ &= -10 \times 1.6990 = -16.99. \end{aligned}$$

Reviewing the earlier lessons, power in watts equals E^2/R or I^2R . Therefore, the equation can be written:

$$\text{db} = 10 \log \frac{(I^2R)_2}{(I^2R)_1}$$

APPENDIX B (Continued)

and if the resistances or impedances are equal, they cancel, and:

$$\text{dB} = 10 \log \frac{I_2^2}{I_1^2}$$

According to the charts at the end of this lesson, the logarithm of the square of a number is equal to two times the logarithm of the number. For example:

$$\log 3 = .4771$$

$$3^2 = 9 \text{ and } \log 9 = .9542.$$

As this relationship holds true for all numbers, the equation for current or voltage ratios is written as:

$$\text{dB} = 10 \log \left(\frac{I_2}{I_1} \right)^2$$

$$= 10 \times 2 \log \frac{I_2}{I_1}$$

$$\text{dB} = 20 \log \frac{I_2}{I_1}$$

$$\text{dB} = 20 \log \frac{E_2}{E_1}$$

Some problems make it necessary to convert a logarithm to its corresponding number. To identify this procedure, it is customary to use the term ANTI-LOG. Thus,

$$\log 3 = .4771$$

$$\text{antilog } .4771 = 3.$$

To illustrate this type of problem, assume that it is desired to calculate the input power of an amplifier which has a gain of 25 dB and an output of 30 watts. Substituting in the equation:

$$\text{dB} = 10 \log \frac{P_2}{P_1}, \quad 25 = 10 \log \frac{30}{P_1}$$

APPENDIX B (Continued)

Dividing both sides of the equation by 10:

$$2.5 = \log \frac{30}{P_1}$$

Taking the antilog of both sides:

$$\text{antilog } 2.5 = \frac{30}{P_1}$$

From the chart, the mantissa ".5" is the log of 3.16 approximately, and the characteristic "2" moves the decimal point two places to the right. Thus,

$$316 = \frac{30}{P_1}$$

$$P_1 = \frac{30}{316} = .095 \text{ watt.}$$

As another example, assume a three stage amplifier, which, with a voltage gain of 20 per stage, has an overall gain of $20 \times 20 \times 20 = 8000$. Expressed in decibels, the gain for each stage is:

$$\begin{aligned} \text{dB} &= 20 \log \frac{E_2}{E_1} = 20 (\log 20) \\ &= 20 \times 1.301, \text{ or } 26.02. \end{aligned}$$

For the overall voltage gain:

$$\begin{aligned} \text{dB} &= 20 \log \frac{E_2}{E_1} = 20 \log 8000 \\ &= 20 \times 3.903, \text{ or } 78.06 \end{aligned}$$

which is a sum of the gains for each stage. That is:

$$26.02 + 26.02 + 26.02 = 78.06 \text{ dB.}$$

EXPRESSED IN dB, THE TOTAL AMPLIFIER GAIN IS EQUAL TO THE SUM OF THE STAGE GAINS.

APPENDIX B (Continued)

Because the decibel represents a ratio between two power values, it is convenient to establish an arbitrary reference level which can be considered as 0 dB. One common reference level for power is .006 watt which is 6 milliwatts. When voltage and current ratios are used, the common reference resistance is either 600 ohms or 500 ohms.

On this basis, for an amplifier which, with an input of 6 milliwatts, provides an output of 24 watts:

$$\text{Power gain} = \frac{P_2}{P_1} = \frac{24}{.006} = 4000$$

$$\text{dB} = 10 \log \frac{P_2}{P_1} = 10 \log 4000$$

$$= 10 \times 3.602 = 36.02.$$

CHART 1

N	0	1	2	3	4	5	6	7	8	9
1.0	.0000	.0043	.0086	.0128	.0170	.0212	.0253	.0294	.0334	.0374
1.1	.0414	.0453	.0492	.0531	.0569	.0607	.0645	.0682	.0719	.0755
1.2	.0792	.0828	.0864	.0899	.0934	.0969	.1004	.1038	.1072	.1106
1.3	.1139	.1173	.1206	.1239	.1271	.1303	.1335	.1367	.1399	.1430
1.4	.1461	.1492	.1523	.1553	.1584	.1614	.1644	.1673	.1703	.1732
1.5	.1761	.1790	.1818	.1847	.1875	.1903	.1931	.1959	.1987	.2014
1.6	.2041	.2068	.2095	.2122	.2148	.2175	.2201	.2227	.2253	.2279
1.7	.2304	.2330	.2355	.2380	.2405	.2430	.2455	.2480	.2504	.2529
1.8	.2553	.2577	.2601	.2625	.2648	.2672	.2695	.2718	.2742	.2765
1.9	.2788	.2810	.2833	.2856	.2878	.2900	.2923	.2945	.2967	.2989
2.0	.3010	.3032	.3054	.3075	.3096	.3118	.3139	.3160	.3181	.3201
2.1	.3222	.3243	.3263	.3284	.3304	.3324	.3345	.3365	.3385	.3404
2.2	.3424	.3444	.3464	.3483	.3502	.3522	.3541	.3560	.3579	.3598
2.3	.3617	.3636	.3655	.3674	.3692	.3711	.3729	.3747	.3766	.3784
2.4	.3802	.3820	.3838	.3856	.3874	.3892	.3909	.3927	.3945	.3962
2.5	.3979	.3997	.4014	.4031	.4048	.4065	.4082	.4099	.4116	.4133
2.6	.4150	.4166	.4183	.4200	.4216	.4232	.4249	.4265	.4281	.4298
2.7	.4314	.4330	.4346	.4362	.4378	.4393	.4409	.4425	.4440	.4456
2.8	.4472	.4487	.4502	.4518	.4533	.4548	.4564	.4579	.4594	.4609
2.9	.4624	.4639	.4654	.4669	.4683	.4698	.4713	.4728	.4742	.4757
3.0	.4771	.4786	.4800	.4814	.4829	.4843	.4857	.4871	.4886	.4900
3.1	.4914	.4928	.4942	.4955	.4969	.4983	.4997	.5011	.5024	.5038
3.2	.5051	.5065	.5079	.5092	.5105	.5119	.5132	.5145	.5159	.5172
3.3	.5185	.5198	.5211	.5224	.5237	.5250	.5263	.5276	.5289	.5302
3.4	.5315	.5328	.5340	.5353	.5366	.5378	.5391	.5403	.5416	.5428
3.5	.5441	.5453	.5465	.5478	.5490	.5502	.5514	.5527	.5539	.5551
3.6	.5563	.5575	.5587	.5599	.5611	.5623	.5635	.5647	.5658	.5670
3.7	.5682	.5694	.5705	.5717	.5729	.5740	.5752	.5763	.5775	.5786
3.8	.5798	.5809	.5821	.5832	.5843	.5855	.5866	.5877	.5888	.5899
3.9	.5911	.5922	.5933	.5944	.5955	.5966	.5977	.5988	.5999	.6010
4.0	.6021	.6031	.6042	.6053	.6064	.6075	.6085	.6096	.6107	.6117
4.1	.6128	.6138	.6149	.6160	.6170	.6180	.6191	.6201	.6212	.6222
4.2	.6232	.6243	.6253	.6263	.6274	.6284	.6294	.6304	.6314	.6325
4.3	.6335	.6345	.6355	.6365	.6375	.6385	.6395	.6405	.6415	.6425
4.4	.6435	.6444	.6454	.6464	.6474	.6484	.6493	.6503	.6513	.6522
4.5	.6532	.6542	.6551	.6561	.6571	.6580	.6590	.6599	.6609	.6618
4.6	.6628	.6637	.6646	.6656	.6665	.6675	.6684	.6693	.6702	.6712
4.7	.6721	.6730	.6739	.6749	.6758	.6767	.6776	.6785	.6794	.6803
4.8	.6812	.6821	.6830	.6839	.6848	.6857	.6866	.6875	.6884	.6893
4.9	.6902	.6911	.6920	.6928	.6937	.6946	.6955	.6964	.6972	.6981
5.0	.6990	.6998	.7007	.7016	.7024	.7033	.7042	.7050	.7059	.7067
5.1	.7076	.7084	.7093	.7101	.7110	.7118	.7126	.7135	.7143	.7152
5.2	.7160	.7168	.7177	.7185	.7193	.7202	.7210	.7218	.7226	.7235
5.3	.7243	.7251	.7259	.7267	.7275	.7284	.7292	.7300	.7308	.7316
5.4	.7324	.7332	.7340	.7348	.7356	.7364	.7372	.7380	.7388	.7396

MANTISSAS OF COMMON LOGARITHMS

CHART 2

N	0	1	2	3	4	5	6	7	8	9
5.5	.7404	.7412	.7419	.7427	.7435	.7443	.7451	.7459	.7466	.7474
5.6	.7482	.7490	.7497	.7505	.7513	.7520	.7528	.7536	.7543	.7551
5.7	.7559	.7566	.7574	.7582	.7589	.7597	.7604	.7612	.7619	.7627
5.8	.7634	.7642	.7649	.7657	.7664	.7672	.7679	.7686	.7694	.7701
5.9	.7709	.7716	.7723	.7731	.7738	.7745	.7752	.7760	.7767	.7774
6.0	.7782	.7789	.7796	.7803	.7810	.7818	.7825	.7832	.7839	.7846
6.1	.7853	.7860	.7868	.7875	.7882	.7889	.7896	.7903	.7910	.7917
6.2	.7924	.7931	.7938	.7945	.7952	.7959	.7966	.7973	.7980	.7987
6.3	.7993	.8000	.8007	.8014	.8021	.8028	.8035	.8041	.8048	.8055
6.4	.8062	.8069	.8075	.8082	.8089	.8096	.8102	.8109	.8116	.8122
6.5	.8129	.8136	.8142	.8149	.8156	.8162	.8169	.8176	.8182	.8189
6.6	.8195	.8202	.8209	.8215	.8222	.8228	.8235	.8241	.8248	.8254
6.7	.8261	.8267	.8274	.8280	.8287	.8293	.8299	.8306	.8312	.8319
6.8	.8325	.8331	.8338	.8344	.8351	.8357	.8363	.8370	.8376	.8382
6.9	.8388	.8395	.8401	.8407	.8414	.8420	.8426	.8432	.8439	.8445
7.0	.8451	.8457	.8463	.8470	.8476	.8482	.8488	.8494	.8500	.8506
7.1	.8513	.8519	.8525	.8531	.8537	.8543	.8549	.8555	.8561	.8567
7.2	.8573	.8579	.8585	.8591	.8597	.8603	.8609	.8615	.8621	.8627
7.3	.8633	.8639	.8645	.8651	.8657	.8663	.8669	.8675	.8681	.8686
7.4	.8692	.8698	.8704	.8710	.8716	.8722	.8727	.8733	.8739	.8745
7.5	.8751	.8756	.8762	.8768	.8774	.8779	.8785	.8791	.8797	.8802
7.6	.8808	.8814	.8820	.8825	.8831	.8837	.8842	.8848	.8854	.8859
7.7	.8865	.8871	.8876	.8882	.8887	.8893	.8899	.8904	.8910	.8915
7.8	.8921	.8927	.8932	.8938	.8943	.8949	.8954	.8960	.8965	.8971
7.9	.8976	.8982	.8987	.8993	.8998	.9004	.9009	.9015	.9020	.9025
8.0	.9031	.9036	.9042	.9047	.9053	.9058	.9063	.9069	.9074	.9079
8.1	.9085	.9090	.9096	.9101	.9106	.9112	.9117	.9122	.9128	.9133
8.2	.9138	.9143	.9149	.9154	.9159	.9165	.9170	.9175	.9180	.9186
8.3	.9191	.9196	.9201	.9206	.9212	.9217	.9222	.9227	.9232	.9238
8.4	.9243	.9248	.9253	.9258	.9263	.9269	.9274	.9279	.9284	.9289
8.5	.9294	.9299	.9304	.9309	.9315	.9320	.9325	.9330	.9335	.9340
8.6	.9345	.9350	.9355	.9360	.9365	.9370	.9375	.9380	.9385	.9390
8.7	.9395	.9400	.9405	.9410	.9415	.9420	.9425	.9430	.9435	.9440
8.8	.9445	.9450	.9455	.9460	.9465	.9469	.9474	.9479	.9484	.9489
8.9	.9494	.9499	.9504	.9509	.9513	.9518	.9523	.9528	.9533	.9538
9.0	.9542	.9547	.9552	.9557	.9562	.9566	.9571	.9576	.9581	.9586
9.1	.9590	.9595	.9600	.9605	.9609	.9614	.9619	.9624	.9628	.9633
9.2	.9638	.9643	.9647	.9652	.9657	.9661	.9666	.9671	.9675	.9680
9.3	.9685	.9689	.9694	.9699	.9703	.9708	.9713	.9717	.9722	.9727
9.4	.9731	.9736	.9741	.9745	.9750	.9754	.9759	.9763	.9768	.9773
9.5	.9777	.9782	.9786	.9791	.9795	.9800	.9805	.9809	.9814	.9818
9.6	.9823	.9827	.9832	.9836	.9841	.9845	.9850	.9854	.9859	.9863
9.7	.9868	.9872	.9877	.9881	.9886	.9890	.9894	.9899	.9903	.9908
9.8	.9912	.9917	.9921	.9926	.9930	.9934	.9939	.9943	.9948	.9952
9.9	.9956	.9961	.9965	.9969	.9974	.9978	.9983	.9987	.9991	.9996

MANTISSAS OF COMMON LOGARITHMS

IMPORTANT DEFINITIONS

ACCELEROMETER — A transducer for measuring the rate of change of velocity, or acceleration.

CARBON MICROPHONE — A cup-shaped housing filled with carbon granules. The carbon microphone is a variable-resistance transducer.

DYNAMIC TRANSDUCER — A form of transducer with a movable coil mounted in a magnetic field. When a mechanical force causes the coil to vibrate, a voltage is induced in the coil. When a varying voltage is impressed across the coil, the coil vibrates mechanically.

INPUT TRANSDUCER — Any form of energy converter which will produce electric energy when some other type of energy is applied to it.

LOUDSPEAKER — A device which converts electric energy to sound energy.

MAGNETOSTRICTION — The property of an iron or steel rod to change in length due to changes of the magnetic flux it carries. Also, the property of a permanent magnet to change its magnetic flux when it is compressed or elongated.

MICROPHONE — A device which converts sound energy to electric energy.

MICROPHONE TRANSFORMER — A device used to match impedances between a microphone output and an amplifier input through transformer action.

OUTPUT TRANSDUCER — Any form of energy converter which will produce another form of energy when electric energy is applied to it.

OUTPUT TRANSFORMER — A device used to match impedances between an amplifier output and a speaker voice coil through transformer action.

PIEZOELECTRIC EFFECT — The property of some crystals to generate a voltage when compressed by a mechanical force and to vibrate mechanically when an alternating voltage is impressed across them.

TRANSDUCER — Any device which converts energy from one form to another with an output dependent on the input.

ULTRASONICS — The science of the production, reception, detection, and absorption of acoustic waves at frequencies of 20 kHz and above.

IMPORTANT DEFINITIONS (Continued)

VARIABLE CAPACITANCE TRANSDUCER — An input transducer in which the input energy causes corresponding changes of circuit capacitance. The changing capacitance is used to produce changing current variations through a load resistor.

VARIABLE RELUCTANCE TRANSDUCER — A form of input transducer in which the input energy causes corresponding changes in the reluctance of a magnetic circuit.

VARIABLE RESISTANCE TRANSDUCER — Any form of input transducer in which the input energy causes corresponding changes in its electrical resistance.

PRACTICE EXERCISE SOLUTIONS

1. An input transducer converts nonelectric energy to electric energy, whereas an output transducer converts electric energy to nonelectric energy.
2. the microphone.
3. True — No energy conversion takes place in the electronic system.
4. (b) an input transducer.
5. (a) both electric. — In general, it is not the purpose of the audio amplifier or electronic systems to convert energy from one form to another. They serve only to amplify a signal, convert it to another frequency, or change its shape.
6. (a) an output transducer. — Electric energy is converted to nonelectric energy.
7. (c) resistance.
8. (a) an input transducer. — The variable resistor converts mechanical motion or displacement to a change of electric resistance.
9. The electric meter. It converts changes of electric current to a mechanical displacement of the pointer on the meter face.
10. True
11. (c) a variable resistor.
12. False — The carbon microphone is not a voltage generator in itself. When placed in an electric circuit, the microphone only varies the circuit resistance. The electric energy required to produce an output is supplied by a battery or other dc source.
13. carbon microphone.
14. (b) reluctance.
15. As the air gap is increased and decreased, the reluctance of the magnetic circuit also increases and decreases. This causes the magnetic field about the circuit to vary. As a result, magnetic lines of force cut the coil, inducing a voltage.
16. The side-to-side movement of the stylus and armature of a variable reluctance cartridge is caused by the variations of the record groove walls.

PRACTICE EXERCISE SOLUTIONS (Continued)

17. (b) free end.
18. (a) increases.
19. (b) aiding.
20. True — Like the carbon microphone, the capacitor microphone does not generate a voltage. It only varies the circuit capacitance. Its power is supplied by an external voltage source.
21. Voltage E_s of Figure 8 is impressed across the diaphragm and fixed plate of the microphone through resistor R_L . As the diaphragm vibrates, the circuit capacitance varies, causing the circuit current to vary accordingly. This changing current flows through R_L to develop the output signal.
22. False — Mechanically moving a coil within a magnetic field will induce a voltage into the coil. Also, applying an ac voltage to the coil will cause the coil to vibrate or move mechanically.
23. Dynamic loudspeakers are classified as permanent magnet (PM) and electromagnetic (EM).
24. (a) an output transducer. — The loudspeaker converts electric energy to sound energy.
25. (b) an electromagnet.
26. (a) alternating current passes through it — Direct current would produce a steady magnetic field like that of an electromagnet, and would not cause the voice coil and cone to vibrate.
27. dynamic
28. True — The dynamic microphone of a public address system converts sound energy to electric energy, whereas the loudspeaker of the same system converts electric energy to sound energy.
29. False — A dynamic microphone generates its own voltage and does not require an external voltage source.
30. piezoelectric
31. True
32. The crystal vibrates mechanically at the frequency of the applied ac voltage.

PRACTICE EXERCISE SOLUTIONS (Continued)

- 33. (b) ac. — When the magnetic bar within the coil of wire is compressed, then elongated, an ac voltage is induced in the coil.
- 34. resonant
- 35. False — Magnetostrictive transducers can also convert electric energy to mechanical energy.
- 36. False — The widest application of magnetostrictive transducers is in the field of ultrasonics.
- 37. True
- 38. False — Maximum transfer of power occurs when the load impedance is equal to the source impedance.
- 39. False — The impedance of a dynamic microphone sees the impedance of the primary as a matching impedance, and not the secondary impedance.



ONE OF THE

BELL & HOWELL SCHOOLS

1082A

1082A EXAMINATION CHECK SHEET

1. 1. A - ALL TRANSDUCERS CONVERT -- ENERGY FROM ONE FORM TO ANOTHER.
Since various forms of energy are converted by transducers, the above general statement best describes transducer action.
2. 2. D - THE CARBON MICROPHONE IS A -- VARIABLE RESISTANCE TRANSDUCER.
Sound pressure waves cause proportional changes in the contact area between loosely packed carbon granules, thus varying circuit resistance.
3. 3. D - TO PRODUCE THE PROPER OUTPUT VOLTAGE, THE ARMATURE OF THE VARIABLE RELUCTANCE PHONOGRAPH CARTRIDGE OF FIGURE 6B MUST MOVE -- FROM SIDE TO SIDE.
Movement of the stylus in a monophonic record groove is side-to-side, and the pole pieces are mounted in such a way that horizontal motion of the armature produces a variable reluctance in the magnetic circuit.
4. 4. B - THE CAPACITANCE OF THE MICROPHONE IN FIGURE 8 IS VARIED BY CHANGING -- THE DISTANCE BETWEEN THE PLATES.
A diaphragm and fixed metal plate function as the plates of an air dielectric capacitor, the distance between which is varied by sound wave pressures.
5. 5. D - THE PURPOSE OF THE FIELD COIL AROUND THE POLE PIECE OF AN EM LOUDSPEAKER IS -- TO PRODUCE A STEADY MAGNETIC FIELD AROUND THE COIL AND POLE PIECE, THUS FORMING AN ELECTROMAGNET.
The field coil and pole piece form the electromagnet in an EM loudspeaker which magnetizes the pole piece. The electromagnet serves the function of the permanent magnet in a PM loudspeaker system.
6. 6. D - A DYNAMIC MICROPHONE WORKS ON THE SAME PRINCIPLE AS A -- DYNAMIC LOUDSPEAKER.
In the dynamic microphone, a pressure diaphragm and attached coil move in a magnetic field. In the loudspeaker, the voice coil attached to a cone moves in a magnetic field.
7. 7. D - A COMMON TYPE PHONOGRAPH PICKUP WHICH WORKS ON THE PIEZOELECTRIC EFFECT IS THE -- CRYSTAL PICKUP.
The transducer element in a cartridge of this type is a crystal such as quartz, tourmaline, or rochelle salt which exhibits the property of piezoelectricity.
8. 8. D - TRANSDUCERS COMMONLY FOUND IN THE FIELD OF ULTRASONICS ARE -- MAGNETO-STRICTIVE TRANSDUCERS.
Magnetostrictive elements are used since high energy ac levels are converted to high energy mechanical vibrations, and other output transducers are not capable of producing or withstanding the required levels of energy transfer.
9. 9. A - THE REASON FOR USING A TRANSFORMER TO COUPLE A DYNAMIC MICROPHONE TO AN AUDIO AMPLIFIER IS TO -- OBTAIN THE BEST PERFORMANCE FROM THE TRANSDUCER BY ENSURING MAXIMUM POWER TRANSFER.
A microphone transformer provides impedance matching between the low impedance microphone output and the high impedance amplifier input, thus providing maximum power transfer.
10. 10. D - IN FIGURE 18, THE SECONDARY WINDING OF THE OUTPUT TRANSFORMER MATCHES THE IMPEDANCE OF -- THE VOICE COIL.
The voice coil is connected to the secondary winding of the output transformer, and "sees" the winding as a matching impedance.

PRACTICE EXERCISE SOLUTIONS (Continued)

- 33. (b) ac. —** When the magnetic bar within the coil of wire is compressed, then elongated, an ac voltage is induced in the coil.
- 34. resonant**
- 35. False —** Magnetostrictive transducers can also convert electric energy to mechanical energy.
- 36. False —** The widest application of magnetostrictive transducers is in the field of ultrasonics.
- 37. True**
- 38. False —** Maximum transfer of power occurs when the load impedance is equal to the source impedance.
- 39. False —** The impedance of a dynamic microphone sees the impedance of the primary as a matching impedance, and not the secondary impedance.

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1082A

Example: Most cars run on

A	
B	
C	
D	

A. coal. B. gasoline. C. ether. D. turpentine.

1. A ☒ **All transducers convert**
 B ☐
 C ☐
 D ☐
 (A) energy from one form to another. (B) electricity to light. (C) electric energy to heat energy. (D) nonelectric energy to electric energy.

2. A ☐ **The CARBON microphone is a**
 B ☐
 C ☐
 D ☒
 (A) dynamic transducer. (B) variable reluctance transducer. (C) variable capacitance transducer. (D) variable resistance transducer.

3. A ☐ **To produce the proper output voltage, the armature of the variable reluctance phonograph cartridge of Figure 6B must move**
 B ☐
 C ☐
 D ☒
 (A) up and down. (B) in a circle. (C) backward and forward. (D) from side to side.

4. A ☐ **The CAPACITANCE of the microphone in Figure 8 is varied by changing**
 B ☒
 C ☐
 D ☐
 (A) the type of dielectric between the plates. (B) the distance between the plates. (C) the resistance of R_L . (D) the area of the plates.

5. A ☐ **The purpose of the field coil around the pole piece of an EM loudspeaker is**
 B ☐
 C ☐
 D ☒
 (A) to transfer signals from the output transformer to the voice coil. (B) to form a complete magnetic path for the pole piece. (C) to alternately magnetize and demagnetize the pole piece in proportion to the signal fluctuations. (D) to produce a steady magnetic field around the coil and pole piece, thus forming an electro-magnet.

6. A ☐ **A DYNAMIC microphone works on the same principle as a**
 B ☐
 C ☐
 D ☒
 (A) carbon pile. (B) potentiometer. (C) variable reluctance cartridge. (D) dynamic loudspeaker.

7. A ☐ **A common type of phonograph pickup which works on the piezoelectric effect is the**
 B ☐
 C ☐
 D ☒
 (A) variable reluctance pickup. (B) variable resistance pickup. (C) dynamic pickup. (D) crystal pickup.

8. A ☐ **Transducers commonly found in the field of ultrasonics are**
 B ☐
 C ☐
 D ☒
 (A) variable resistance transducers. (B) loudspeakers. (C) carbon microphones. (D) magnetostrictive transducers.

9. A ☐ **The reason for using a transformer to couple a dynamic microphone to an audio amplifier is to**
 B ☐
 C ☐
 D ☐
 (A) obtain the best performance from the transducer by ensuring maximum power transfer. (B) to convert the transducer's output energy into electrical energy. (C) step up the transducer's output as high as possible, thus obtaining sufficient power to operate a speaker directly. (D) match a high impedance source to a low impedance load.

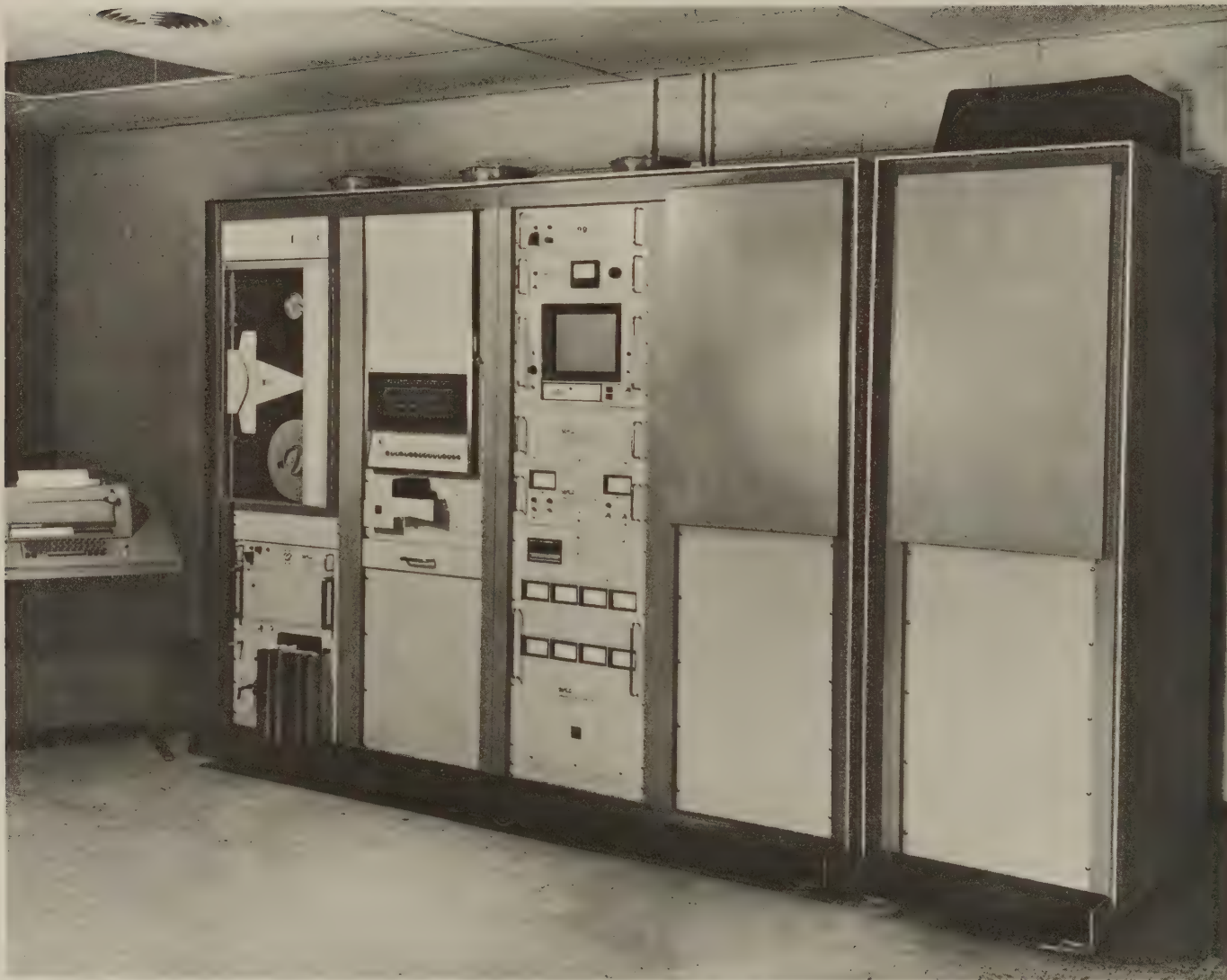
10. A ☐ **In Figure 18, the SECONDARY WINDING of the output transformer matches the impedance of**
 B ☐
 C ☐
 D ☒
 (A) the primary winding. (B) the amplifier output. (C) the speaker. (D) the voice coil.

ELECTRONIC SYSTEMS

1085

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





Electronic systems may range in size from tiny hearing aid units weighing less than 1 ounce, to complex intercontinental military and space systems. An almost limitless number of electronic system types can be assembled from smaller subsystems and standard electronic components. Computers are often an important part of large systems such as this automated artwork generator for integrated circuits.

Courtesy Seaco Computer Display Corp.

CONTENTS

Circuit Functions	Page 4
Analog Circuits	Page 5
Digital Circuits	Page 5
Amplifiers	Page 7
Distortion and Fidelity	Page 7
Gain	Page 8
Bandwidth	Page 8
Oscillators	Page 10
Communication Systems	Page 11
Radio Systems	Page 11
Transmitters	Page 12
Modulation	Page 14
Receivers	Page 16
Television Systems	Page 18
Radar Systems	Page 22
Industrial Electronic Systems	Page 25
Information Systems	Page 26
Computer Systems	Page 27
Information Storage and Retrieval Systems	Page 29

Knowledge once gained casts a faint light beyond its own immediate boundaries. There is no discovery so limited as not to illuminate something beyond itself.

—John Tyndall

ELECTRONIC SYSTEMS

An electronic system is defined as a group of components interconnected to perform a function or group of related functions. If a number of systems are interconnected into a larger system, the individual sections are then called subsystems.

The age of electronics has also brought the age of high-speed communication and information exchange. Most of the electronic equipment in use today is part of some COMMUNICATION or INFORMATION system. This lesson describes several of the wide variety of systems used in homes, industries, and military installations.

There are two general types of systems: **ANALOG** and **DIGITAL**. Analog systems are most often used for communication between people, while digital systems are primarily used to provide communication between machines, or between man and machine. For the purpose of this lesson, we may say that analog systems use **ANALOG CIRCUITS**, and digital systems use **DIGITAL CIRCUITS**. Systems which use both analog and digital circuits are called **HYBRID SYSTEMS**.

CIRCUIT FUNCTIONS

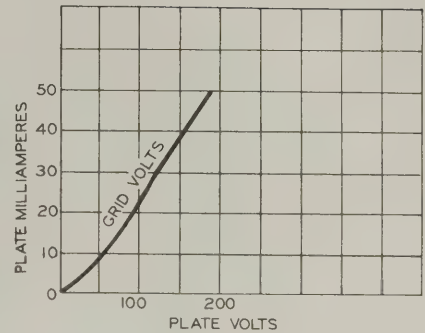
All of the complex functions performed by electronic circuits can be reduced to the functions of **AMPLIFICATION**, **OSCILLATION** and **SWITCHING**. Circuits that perform amplification or oscillation are **ANALOG CIRCUITS**. Those that perform switching only are **DIGITAL CIRCUITS**. All electronic equipment uses these basic circuit types in various combinations to function within a wide variety of electronic systems.

Analog circuits exhibit unique properties in that they will **AMPLIFY** and **OSCILLATE**. These two functions are so important that they actually provide the basis for electronics as we know it. Prior to the invention of the vacuum tube, electrical circuits could only be **SWITCHED**, as was done in early telephone and telegraph systems. Vacuum tubes and transistors have added the dimensions of amplification and oscillation to electrical circuitry, and are called **ACTIVE ELEMENTS**. However, most active elements will not amplify or oscillate by themselves. They require external circuitry to route voltages, provide feedback, and perform a variety of functions. This complementary circuitry, or **PASSIVE ELEMENTS**, consists of resistors, capacitors and inductors. Similarly, passive elements will not amplify or oscillate by themselves. In actual practice, active and passive elements are

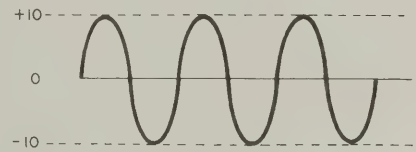
connected together to provide amplification, oscillation and switching functions. These, in turn, are connected together to form analog or digital circuits. These circuits are then connected together to form analog, digital or hybrid systems.

Analog Circuits

Amplifiers and oscillators are analog circuits which are commonly studied in electronics courses, and are the type with which you are probably already most familiar. The words ANALOG and ANALOGOUS are defined by Webster as "... according to a due ratio, proportionate." and "... having analogy; corresponding in some respects to something else ...". An analog circuit, therefore, is one in which the output corresponds in some proportionate way to the input. Stated more precisely, an analog circuit is one in which the output varies smoothly over a given range, and therefore has an infinite number of discrete voltage or current values which correspond to some portion of the input. As will be shown with the aid of Figure 1, amplifiers and oscillators possess this quality. Figure 1A is a typical characteristic curve for an electron tube. The output of the electron tube may vary almost linearly over a range of 0 to 150 volts. In other words, the tube could operate under steady-state conditions at ANY POINT along the curve, depending upon input signal conditions. Thus, an amplifier output is not only a faithful reproduction of the input, it is also proportional to the input, and is therefore an analogous representation of the input.



A



B

Figure 1

Oscillators are also analog circuits since they produce alternating waveforms, such as that shown in Figure 1B, and also because the oscillations are sustained through the use of **FEEDBACK**, which is a true analog function. Analog oscillators always use positive feedback. That is, a portion of the output of an analog oscillator circuit is fed back to the input in phase with the input signal. This means that the circuit is capable of sustaining oscillation since only a small portion of the output is fed back. The oscillator stage is also an amplifier, and therefore, the output is also an amplified reproduction of the input. With only a portion of the output being fed back, the original external input signal can be removed from the circuit, and the oscillator will continue to operate on its own.

Digital Circuits

A digital circuit is one in which an input signal can produce one of only two separate and distinct output states. This may be thought of pictorially as a simple switch where the device is either "on" or "off". Note that in Figure 2, when the switch is closed, a circuit is completed from the battery through the lamp. In this circuit, the lamp can be either on or off. The only operating voltage available in Figure 2 is the 12V dc battery supply, and no intermediate operating level can be selected. The switching element is, in

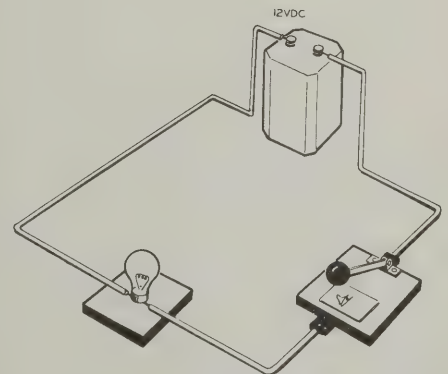
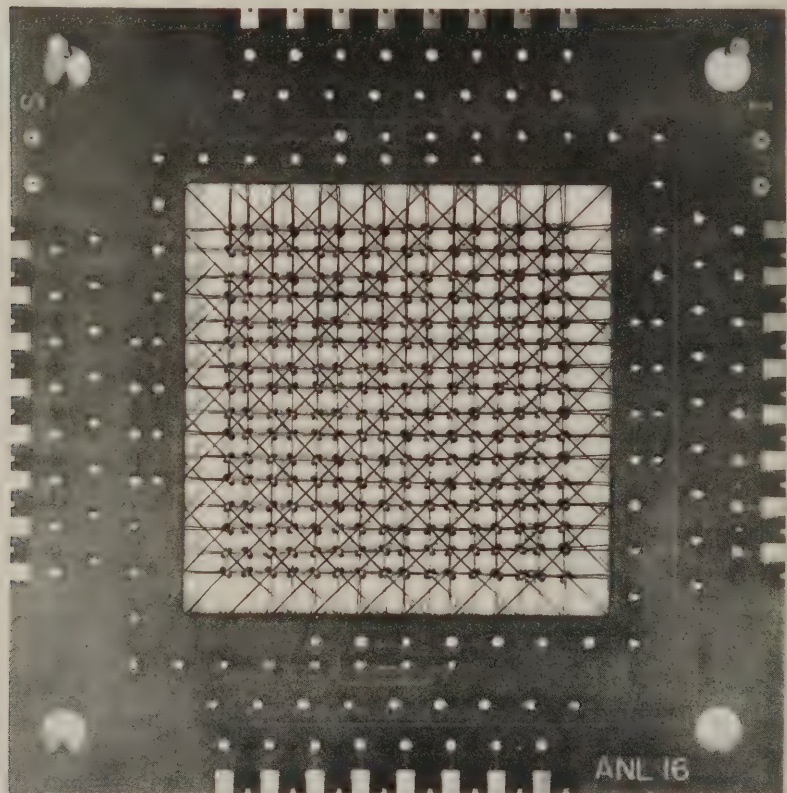


Figure 2

this example, a mechanical switch. However, a variety of electronic components are capable of simple on-off switching. Such devices can be connected as **BISTABLE ELEMENTS**, and include relays, diodes and ferromagnetic cores used in computer memories. Electron tubes and transistors can also be used in circuit configurations which permit them to be only on or off. Connected in this way, tubes and transistors become **SWITCHING** function **AMPLIFIERS** rather than analog amplifiers or oscillators, and instead are part of a digital circuit. In one such configuration, a tube or transistor circuit is either conducting, or it is cut off, and is referred to as a **FLIP-FLOP** circuit. These circuits are so named because an input pulse causes the device to “flip” or “flop” between two steady-output states. In some cases, however, flip-flop circuits may be allowed to “vibrate” rather than remain in one steady state. These circuits are actually astable multivibrators, and although the output is periodic, these devices are performing a digital circuit function and in this respect are different from analog oscillators. A bistable multivibrator used in a digital system may be thought of as “on” when in the vibrating, or periodic state, and “off” when the device is not vibrating. Seen in this way, the multivibrator actually has only two



Ferrite cores, such as the ones strung together on this computer memory frame, are true digital devices. A pulse of electric energy appearing on one of the lines passing through the center of a core causes the core to be magnetically polarized in one of two possible states.

Courtesy Becks, Inc.

states, "on" and "off", and is fulfilling a digital circuit function. In addition, it is noted that the periodic alternations usually have a fixed amplitude and frequency. The only condition that is controlled by the input with such a device is the on-off function.

AMPLIFIERS

An amplifier circuit usually contains a number of passive elements as well as one or more active elements. As shown in Figure 3, to provide amplification, a tube or transistor must receive operating voltages and currents from a power supply.

In an electronic system, each active element and its associated passive elements are called a **STAGE**. A stage may provide amplification, or it may perform some other function.

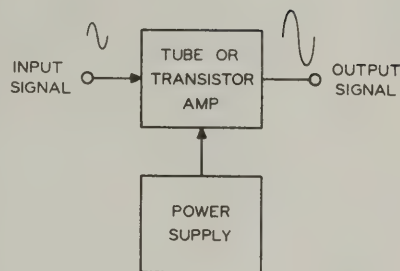


Figure
3

Distortion and Fidelity

It is desirable for the output variations of most amplifiers to resemble the input variations as closely as possible. The output waveform should be an exact replica or copy of the input waveform except for an increase in magnitude. In most single stage amplifiers the output waveform is an **INVERTED** replica of the input waveform. In many applications this inversion in polarity is unimportant as long as the output variations duplicate the input variations. In applications where the inversion would have undesirable effects, however, the stage which produces the undesirable inversion can be followed by another **STAGE** which restores the waveform to its original polarity.

Differences between the output and input waveforms of an amplifier, other than a difference in magnitude and polarity, are referred to as **DISTORTION**. A **HIGH FIDELITY** amplifier is capable of amplifying an input **SIGNAL** with very little distortion. The term fidelity means faithfulness. Thus, the output of a high fidelity amplifier is a **FAITHFUL** reproduction of the input signal.

Fidelity, however, is not always important, especially in some industrial applications. In many industrial applications, the only requirement of an amplifier is that it increase the magnitude of a control signal enough to operate some output device such as an electric motor or relay.

In other applications the amplifier is actually designed to distort the input signal in some way. In FM radio receivers and television sets, for example,

some of the amplifiers may be designed to limit the signal to some constant amplitude.

Gain

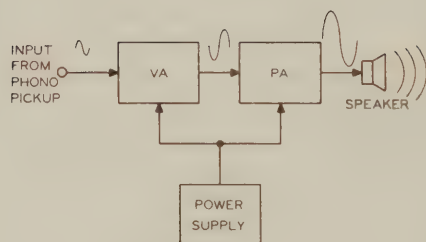


Figure
4

The ratio of the output to the input of an amplifier is called the **AMPLIFICATION FACTOR** or **GAIN** of the amplifier. This ratio is denoted by the symbol A or G .

Depending on the application, voltage gain, current gain or power gain may be of primary interest. For example, an audio amplifier used in a phonograph may contain two stages of amplification, as shown in Figure 4. The main function of the first stage is to provide a voltage gain. This stage is called a voltage amplifier (VA). The main function of the second stage is to provide a power gain. This stage is called a power amplifier (PA).

In a phonograph pickup, the mechanical motion of the stylus, or needle, produces voltage variations. The power amplifier develops the power required to operate the speaker. However, the output signal from the pickup may not have enough amplitude to drive the power amplifier. In this case a voltage amplifier is placed between the pickup and the power amplifier to amplify the signal voltage. Thus, the small voltage variations produced by the pickup are amplified by the voltage amplifier. In turn, the output of the voltage amplifier drives the power amplifier, which develops the power that operates the speaker.

Although an amplifier may provide a power gain, it is important to note that no power is actually created within the amplifier. The increased signal power at the output of an amplifier is obtained from the power supply. In turn, the power supply may derive its power from the ac power line.

Bandwidth

The **BANDWIDTH** of an amplifier is the frequency range over which the amplifier has relatively constant gain. The bandwidth of a voltage amplifier, for example, may be defined as the range of frequencies over which the output voltage is at least 70% of maximum when a constant amplitude input signal is applied. In amplifiers of this type, the output voltage usually decreases rapidly at higher than this range and at frequencies lower than this range. The bandwidth, then, is an indication of the frequency range over which an amplifier output does not fall off more than 3 dB.

Amplifiers are frequently classified according to the range of frequencies over which they are used. Audio amplifiers, for example, are used to amplify audio frequencies ranging from a few hertz (cycles per second) to about 20 kilohertz. In television systems, the frequencies representing the picture signal range from a few hertz to about 4.2 megahertz. Amplifiers used to amplify signals with this wide range of frequencies are called video amplifiers. Radio frequency (r-f) amplifiers are used to amplify signals of radio frequencies. Depending on the application, an r-f amplifier may have anywhere from a very narrow bandwidth to a very wide bandwidth.

Bandwidth may be increased by feeding back, out-of-phase, a portion of the output signal to the input. This use of negative feedback also causes a reduction in distortion through an amplifier.

A system that uses either negative or positive feedback, where a portion of the output is fed back to the input, may be referred to as a **CLOSED LOOP SYSTEM**.

OSCILLATORS

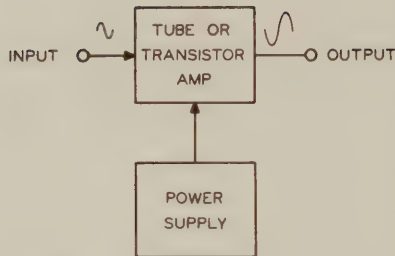


Figure 5

Many electronic systems require one or more alternating voltages or currents. One of the basic requirements for a radio communication system, for example, is a radio frequency voltage or current which is applied to an antenna to produce radio waves. The ac present in the ac power lines cannot be used for most of these applications because the frequencies required are usually much higher than 60 Hz (hertz) and may require waveforms that are not sinusoidal. Some of these waveforms, for example, resemble a sawtooth, a triangle, or a rectangle. These alternating voltages and currents must be produced by circuits within the system itself. These circuits are called **OSCILLATORS**.

An oscillator is made up of an amplifier and a feedback circuit which feeds a portion of the output back to the input of the amplifier. The use of the feedback circuit causes the amplifier to produce its own input once it is operating. The feedback circuit must meet two requirements in order to maintain oscillations. These requirements are:

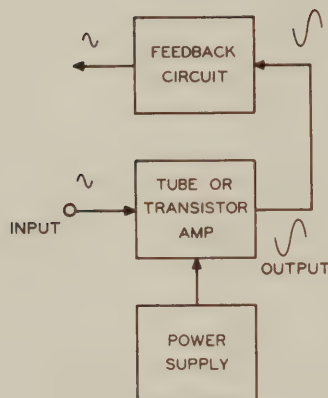


Figure 6

1. The feedback circuit must produce a voltage or current of the proper phase to the input.
2. The feedback circuit must provide a voltage or current of the proper amplitude to the input.

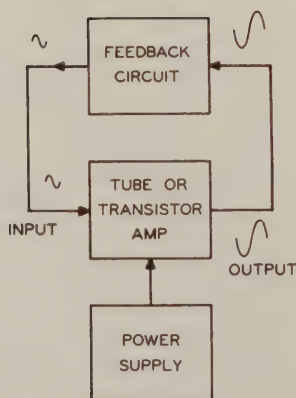


Figure 7

Figures 5, 6 and 7 show the basic principles of a typical oscillator circuit. Figure 5 represents an amplifier in which the output has a greater magnitude than the input and its polarity is reversed. As previously mentioned, this is typical for a single stage of amplification. In Figure 6, the output of the amplifier is applied to a feedback circuit. The feedback circuit reverses the polarity of the output and reduces its amplitude so that it is identical to the amplifier input. Thus, the feedback voltage or current can be used as the amplifier input as shown in Figure 7.

Summarizing the operation of the oscillator, the feedback voltage applied to the amplifier input is amplified and applied to the feedback circuit. The feedback circuit feeds a portion of the amplifier output back to the input and the whole operation repeats itself to maintain oscillations. This action can be started by small changes of voltage or current at the amplifier input when the oscillator is first turned on.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

CIRCUIT FUNCTIONS

ANALOG CIRCUITS

DIGITAL CIRCUITS

AMPLIFIERS

DISTORTION AND FIDELITY

GAIN

BANDWIDTH

1. Vacuum tubes and transistors are _____ elements.
2. Switching circuits can be made using passive elements only. True or False?
3. A typical closed loop system is (a) an oscillator, (b) a digital circuit.
4. Devices which may be in only one of two possible states are (a) active elements, (b) passive elements, (c) bistable elements.
5. Electron tubes and transistors may be used as switches. True or False?
6. The waveform of output voltage or current of an amplifier normally has a greater amplitude than the input waveform and may be inverted. Any other differences between the output and input waveforms are referred to as _____.
7. With reference to electronic amplifiers, what does the term high fidelity mean?
8. A certain amplifier develops an output of 20 volts with an input of 1 volt. What is the voltage gain of the amplifier?
9. With respect to electronic amplifiers, what does the term bandwidth mean?
10. An amplifier which has a relatively constant gain from about 20 hertz to just above 20 kilohertz would be classed as (a) an audio amplifier, (b) a video amplifier.



COMMUNICATION SYSTEMS

The two principal means for rapid communications over long distances are by wire and radio (including TV). Radio is the only practical means for extended distance communications with moving vehicles such as automobiles, boats, aircraft, and spacecraft. Radio communication systems are also more practical for communications over large bodies of water, such as the oceans. Radio communication systems are also more practical where the transmissions are to be received at a great many points, as in broadcasting to the general public.

Since radio communications have greatly influenced the development of electronics, some knowledge of radio communication systems is very helpful in any field of electronics. For example, r-f heaters used in industrial processes are very similar to the transmitters used in radio communication systems. Some of the fields most closely associated with radio communications include radar, radio astronomy, radio control, and telemetering, or measuring at a distance.

Radio Systems

Radio communication systems are used to transmit intelligent information in analog form, and must therefore use analog circuits. It is useful to note here that acoustic waves produced by the human voice are analog in nature. Thus, we find that radio communication systems are ideally suited to man's communication needs. Through radio, human speech and other intelligent sounds may be transmitted from one point to another in natural form, without requiring the use of an alien type of communication. This type of transmission is called **RADIOTELEPHONY**, or simply radio. Radio systems may, however, transmit intelligent information via methods other than those of speech or music.

Another type of radio communication is **RADIOTELEGRAPHY**. In this system, messages are transmitted in code form which must be translated by the person sending the message, and interpreted by the person receiving the message. The code used is a system of dots and dashes known as the International Morse Code. Morse code may be transmitted and received by human operators through the use of a telegraph key. The sender "taps" out the message manually by pressing and releasing the key according to the code. Similarly, the person receiving the message interprets the incoming series of dots and dashes as the original message. This, of course, is time consuming and requires considerable skill on the part of both persons taking part in the communication. Communication in this form has advantages in the electronics aspect of the system since it is the simplest form of radio transmission. However, this system is slower and more difficult for humans to

handle and interpret. Actually, there is a method for speeding up the operation, and one which is in widespread use. If we substitute a form of typewriter for the manual keys, we have a **RADIOTELETYPE** System. This is still a form of radiotelegraphy, however, and is not a new form of transmission. A Radioteletype system, or simply teletype system, merely automates the code generation by forming the appropriate code for each letter when the typewriter key is depressed. The code is then transmitted at high speed via radio waves. The entire code generation system is then reversed at the receiving typewriter (printer), and the message can be read in the form of the printed word, which is easily interpreted and does not require a knowledge of Morse Code.

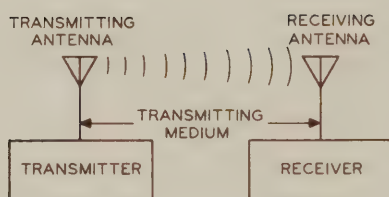


Figure
8

A third type of radio communication is television. Television is actually a radio system in which video information is included in the transmission. This is accomplished through the use of a television camera, which provides visual information for the TV system as the microphone provides audio information in an ordinary radio system. The information is then transmitted via radio waves to a receiver, or home television set. The television set reverses the procedure established by the TV camera, with the picture tube providing the picture compared to the use of a speaker to provide the sound in an ordinary radio system.

The fourth type of radio communication is **RADIOFACSIMILE**, or simply facsimile. This system is used to transmit still pictures via radio waves. Facsimile permits the transmission of "hard copy". That is, an actual picture is obtained from the receiver which may be removed from the machine, handled, and used just as the original photograph is used. Facsimile methods are used to obtain photos from weather satellites, showing cloud cover and other atmospheric detail.

Figure 8 is a block diagram of a radio communication system. The three main parts of this system are the transmitter, the transmitting medium, and the receiver. The transmitter impresses the desired information on r-f voltages and currents which are fed to an antenna. The antenna radiates the radio waves. The transmitting medium is the space and anything in it through which the radio waves propagate. As the radio waves cut the receiving antenna, they induce voltages and currents in it. The receiver extracts the information from the signal picked up by the antenna and reproduces it in usable form.

Transmitters

Figure 9 is a block diagram of a basic radiotelegraphy transmitter. The R-F OSC block represents an electronic oscillator which generates a radio fre-

quency wave. This r-f wave, or CARRIER, is amplified and applied to the transmitting antenna. The resulting radio wave is radiated from the antenna.

A radio frequency carrier transmitted continuously at a constant amplitude and frequency contains no intelligible information. If the radio waves are to carry information to a distant receiver, there must be some way to impress the information on the carrier. The process of impressing information on a carrier wave is called **MODULATION**.

In radiotelegraphy the carrier may be modulated by turning the transmission on and off with a key to form the dots and dashes of the telegraph code. This form of transmission is called continuous wave (cw) telegraphy. As shown in Figure 9, the key is simply a switch. The transmission is turned on and off by depressing and releasing the key. In this simplified example, when the key is depressed power is applied to the oscillator and amplifier and radio waves are radiated from the antenna. When the key is released, as shown in the figure, the oscillator and amplifier are inoperative and no radio waves are radiated from the antenna. Depressing the key for a short period of time forms a "dot" and depressing the key for about three times as long forms a "dash".

Figure 10 is a block diagram of a radiotelephone transmitter. A radiotelephone transmitter is made up of two main sections—a radio frequency (r-f) section and an audio frequency (a-f) section. A power supply provides the required voltages and currents for both sections.

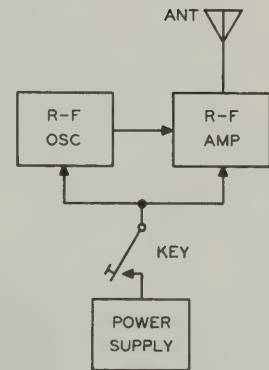


Figure 9

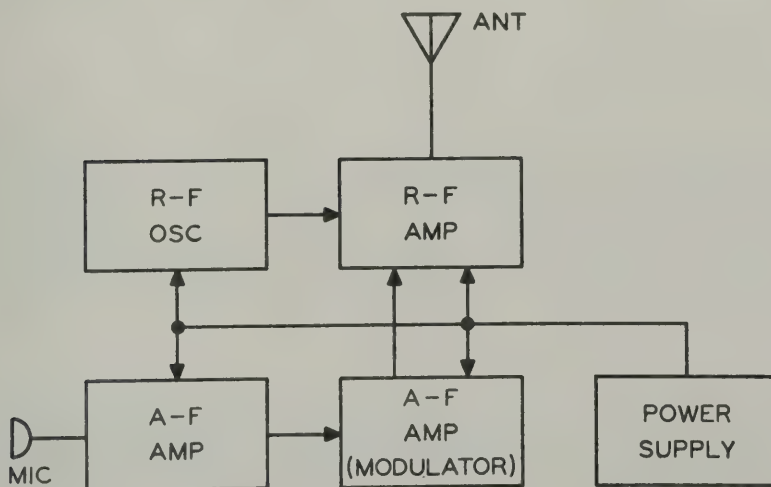


Figure 10

The r-f section is made up of an electronic oscillator which generates the r-f carrier and an r-f amplifier which amplifies the carrier to a suitable level. In this figure, an audio frequency amplifier called a **MODULATOR** applies the a-f signal to the r-f amplifier in such a way as to impress the audio frequency information on the r-f carrier generated by the oscillator. A microphone produces a very weak audio frequency signal. Therefore, another a-f amplifier, in addition to the modulator, is inserted between the microphone and the modulator to obtain the audio frequency power necessary for proper modulation.

Modulation

A transmitter like that of Figure 10 produces a type of modulation known as **AMPLITUDE MODULATION (AM)**. In amplitude modulation, **THE AMPLITUDE OF THE R-F CARRIER VARIES IN ACCORDANCE WITH THE MODULATING SIGNAL**.

FIGURE 11 shows the waveforms of the unmodulated r-f carrier, the modulating a-f signal, and the resulting amplitude modulated r-f carrier. When

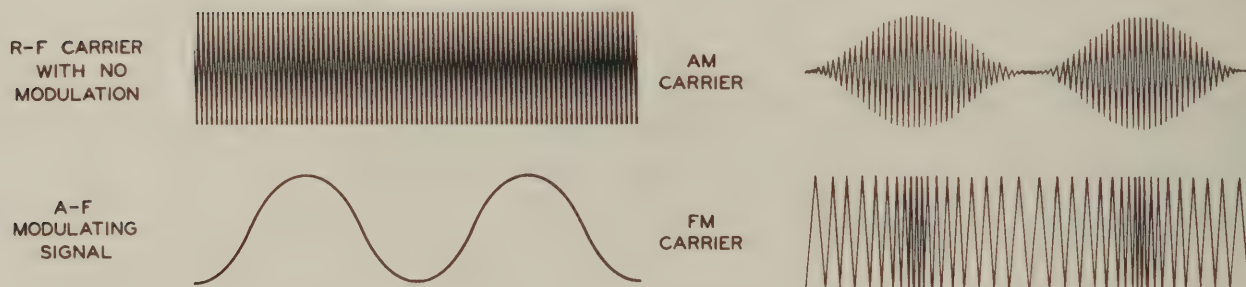


Figure 11

the a-f modulating signal is at its positive peak, the AM carrier has maximum amplitude. When the a-f modulating signal is at its negative peak, the AM carrier has minimum amplitude. The amplitude of the modulating signal determines the amplitude of the modulated carrier. The frequency of the modulating signal determines at what frequency the amplitude of the carrier varies.

Figure 11 also illustrates a type of modulation in which **THE FREQUENCY OF THE R-F CARRIER VARIES IN ACCORDANCE WITH THE MODULATING SIGNAL**. This type of modulation is called **FREQUENCY MODULATION (FM)**.

In the illustration of Figure 11, when the a-f modulating signal is at its positive peak, the frequency of the FM carrier is maximum. When the a-f modulating signal is at its negative peak, the frequency of the FM carrier is minimum. This action may also be reversed so that the maximum and minimum frequencies of the FM carrier correspond to the negative and positive peaks of the modulating signal, respectively. In either case, the frequency of the r-f carrier varies between some maximum and minimum value in accordance with the modulating signal.

In FM systems, the amplitude of the modulating signal determines the amount of SWING or DEVIATION from the unmodulated carrier frequency. The frequency of the modulating signal determines the frequency with which the frequency of the carrier varies. The CENTER FREQUENCY of the modulated carrier is the same as the frequency of the unmodulated carrier.

An FM transmitter is similar to an AM transmitter in many ways. The main difference is in the point at which the modulating signal is impressed on the carrier; that is, the point at which the carrier is modulated. In the AM transmitter of Figure 10, the modulation is applied to the r-f amplifier in such a way as to cause the amplitude of the carrier to vary in accordance with the modulating signal. The stage in which this occurs is usually separate and isolated from the oscillator. In an FM transmitter, the modulator may be connected directly to the r-f oscillator in such a way as to cause the frequency of the oscillations to vary in accordance with the modulating signal.

Receivers

Electromagnetic waves from many different sources induce voltages and currents in the antenna of a radio receiver. Therefore, one of the most important characteristics of a receiver is its ability to select one particular frequency and to reject all others. This characteristic is called **SELECTIVITY**. Furthermore, the desired frequency or r-f carrier may be very weak by the time it reaches the receiving antenna. The ability of a radio receiver to adequately amplify small signals is another important characteristic of radio receivers. This characteristic is called **SENSITIVITY**.

Figure 12 is a block diagram of a type of radio receiver called a **TUNED RADIO FREQUENCY (TRF) RECEIVER**. The first section of the receiver is made up of tuned circuits which select the desired frequency.

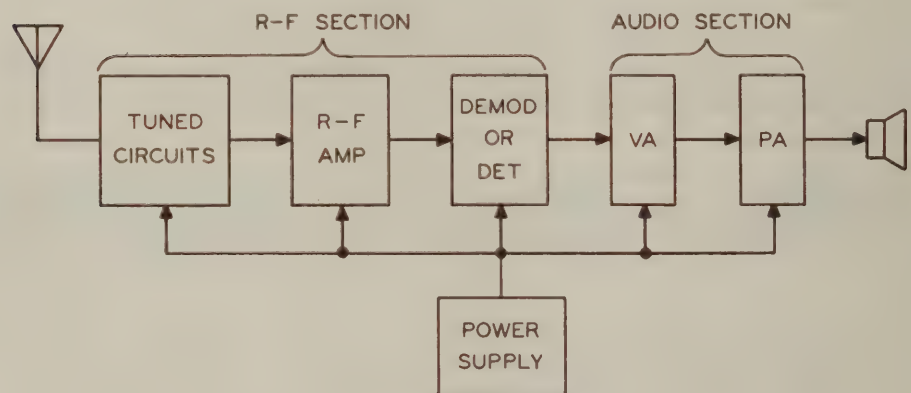


Figure 12

The desired r-f carrier is amplified by the r-f amplifier and applied to a **DEMULATOR**. The demodulator recovers the modulating signal from the modulated r-f carrier. The demodulator is also called a **DETECTOR**. The output from the demodulator is applied to an audio section which may be exactly like the phonograph amplifier described previously. The voltage amplifier (VA) drives a power amplifier (PA) which provides the power to drive the speaker.

In practice, a TRF receiver may contain more than one stage of r-f amplification. Tuned circuits may be inserted between neighboring r-f amplifier stages and between the last r-f amplifier and the demodulator. These tuned circuits all help in selecting the desired frequency. One disadvantage of the TRF receiver is that many tuned circuits may have to be tuned at the same time

The following Practice Exercise questions cover the subjects which you have just studied. They are:

OSCILLATORS

COMMUNICATION SYSTEMS

RADIO SYSTEMS

TRANSMITTERS

MODULATION

11. A circuit which generates an alternating current or voltage using dc from a power supply is called (a) an oscillator, (b) a rectifier.
12. What two circuits are required in an oscillator?
13. What are the two main requirements of the feedback voltage or current to maintain oscillations in an oscillator?
14. The output waveform of an oscillator is always sinusoidal. True or False?
15. Name four types of information commonly transmitted by radio communications systems.
16. What type of circuit is used to develop the r-f carrier in a radio transmitter?
17. The process of impressing information on a carrier wave is called _____.
18. The operation of keying an r-f carrier on and off to form dots and dashes is a form of modulation. True or False?
19. In the radiotelephone transmitter of Figure 10, what type of circuit is used to impress the audio signal on the r-f carrier? (a) a rectifier circuit, (b) an oscillator circuit, (c) an amplifier circuit.
20. In an AM transmitter, what characteristic of the r-f carrier is varied in accordance with the modulating signal?
21. In an FM transmitter, what characteristic of the r-f carrier is varied in accordance with the modulating signal?

to select the desired frequency. Thus, the large number of controls on a TRF receiver may make tuning difficult.

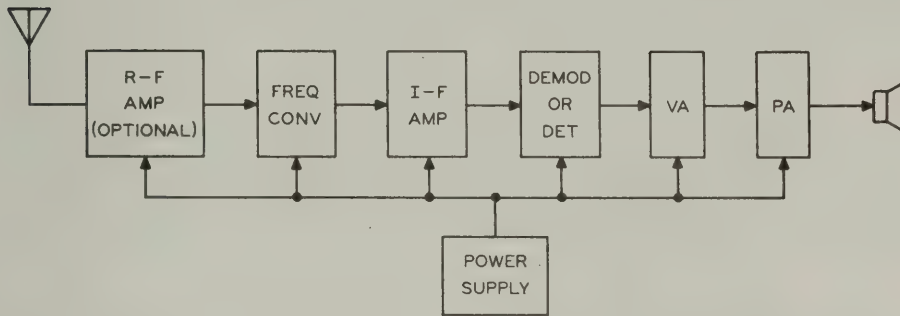


Figure 13

Figure 13 is the block diagram of another type of receiver which requires only one tuning control. This type of receiver is called a **SUPERHETERODYNE RECEIVER**. Almost all home entertainment receivers in use today are of this type. While not shown in this figure, tuned circuits are found in some of the stages of the superheterodyne receiver.

A superheterodyne receiver converts the incoming r-f carrier to a lower frequency called the **INTERMEDIATE FREQUENCY (i-f)**. A frequency converter stage changes all selected r-f signals to the same single intermediate frequency. No matter what r-f carrier frequency is received, the intermediate frequency is always the same.

The amplifiers used to amplify the intermediate frequency are appropriately called i-f amplifiers. These are radio frequency amplifiers, but because they operate at the intermediate radio frequency, the name i-f amplifier is used rather than r-f amplifier. The i-f amplifiers of a superheterodyne receiver are tuned at the factory and then need not be retuned unless some part changes value or is replaced.

The most common intermediate frequency used in AM home receivers is 455 kHz. An i-f of 10.7 MHz is usually used in FM receivers. A constant intermediate frequency eliminates the need to retune the i-f amplifier each time the receiver is tuned to a different r-f carrier frequency.

Now let's summarize the operation of the superheterodyne receiver of Figure 13. First of all, an r-f amplifier may be inserted between the antenna and

the frequency converter. The use of this r-f amplifier, however, is optional and many receivers operate very well without one. The frequency converter, which is made up of a special type of amplifier and an oscillator, converts the desired incoming r-f carrier frequency to the intermediate frequency. The intermediate frequency signal is amplified by one or more i-f amplifiers and then applied to a demodulator.

The intermediate frequency has all the characteristics of the incoming r-f carrier frequency with respect to modulation. The modulating signal is recovered by the demodulator and applied to the audio section which operates the speaker, as previously described.

The main differences between an AM receiver and an FM receiver are the choice of intermediate frequency and the type of demodulator used. As mentioned previously, most home entertainment type AM radios have an i-f of 455 kHz and most FM radios have an i-f of 10.7 MHz. Because of the differences in radio frequency effects at these respective radio frequencies and intermediate frequencies, the wiring and component layout of an FM radio is much more critical than in most AM radios.

Television Systems

The operations of a **TELEVISION** system depends on a special type of tube called a cathode ray tube (CRT). Cathode ray tubes are used as both input and output transducers in the complete television system. A cathode ray tube called a camera tube is used in a television camera to convert light energy to electric energy. The picture to be televised is focused on the camera tube, and the camera tube produces a corresponding electrical signal.

Another type of cathode ray tube, called a picture tube, is used in a television receiver to convert electric energy to light energy. This action is illustrated in Figure 14. The electric signal produced by the camera tube results in a similar signal at the receiver. This signal is applied to the picture tube. The resulting picture produced by the picture tube is the same as that focused on the camera tube.

A section of a cathode ray tube, called an electron gun, produces a narrow beam of electrons which are directed down the length of the tube. Before it was discovered that this beam was made up of electrons, it was called a cathode ray and thus gave rise to the name cathode ray tube. The position of the electron beam is controlled by a magnetic field. A group of coils, called a yoke, is placed around the tube. Passing a current through these coils produces a magnetic field and the magnetic field, in turn, deflects the electron beam.

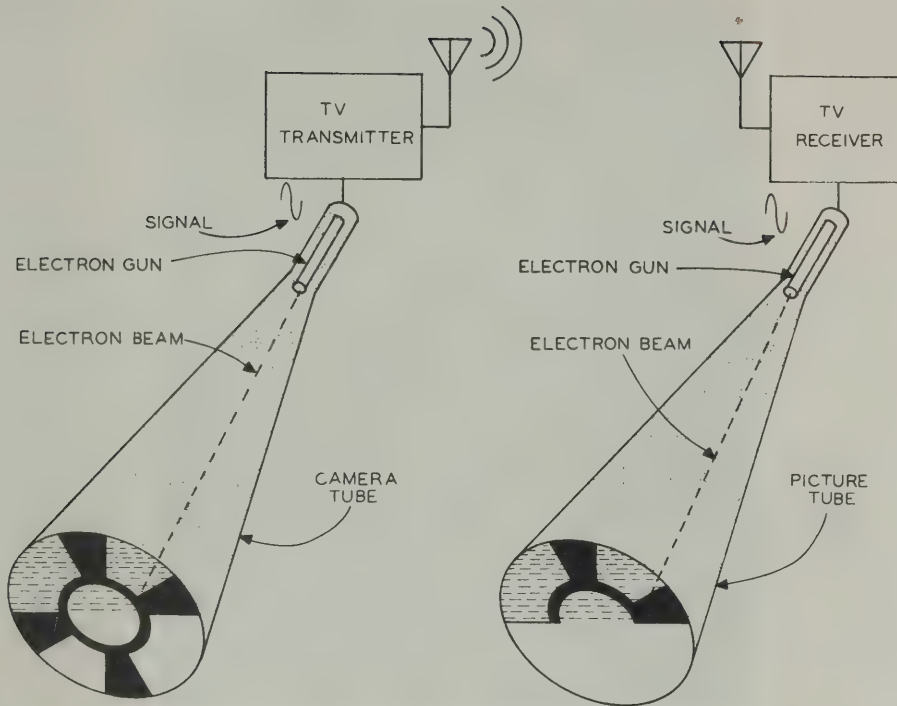


Figure 14

The current passing through the yoke sweeps the beam back and forth across the picture focused on the camera tube. The beam moves downward a little each time until the entire picture has been scanned. The beam is then deflected back to the top of the picture to repeat this sweeping, or scanning, process. As the beam sweeps back and forth across the picture, the camera tube produces an electric signal which corresponds to the light and dark areas of the picture.

The electron beam in the picture tube sweeps the face of the picture tube in the same way that the electron beam in the camera tube sweeps the picture. The face of the picture tube is coated with a fluorescent substance which emits light as it is struck by the electron beam. The amount of light emitted by the fluorescent substance varies with the intensity of the electron beam. Thus, the light and dark areas of the picture are reproduced by varying the intensity of the electron beam.

The intensity of the electron beam is controlled by the electric signal derived from the camera tube. Summarizing then, as the electron beam in the camera tube sweeps the picture, the electron beam in the picture tube "paints" the same picture on the fluorescent screen.



Television systems include a wide variety of electronic equipment. The studio area shown here represents only a small portion of a system which may include network stations and, ultimately, the viewers' home receivers.

Courtesy WAAM-TV, Baltimore, Md.

Figure 15 is a block diagram of a typical television receiver. All television receivers in use today are superheterodyne receivers. The picture and sound carriers are received by the antenna and amplified by the r-f amplifier. The frequency converter converts the modulated picture and sound carriers to modulated i-f signals. Common intermediate frequencies used in television receivers range from about 20 MHz to about 40 MHz.

At some point between the frequency converter and the picture tube, the sound i-f signal is separated from the picture i-f signal. In the television receiver of Figure 15, this separation takes place in the i-f amplifier section. This may be accomplished by resonant circuits tuned to the respective sound and picture i-f's. The sound i-f is fed to the sound section and the picture i-f is fed to the picture section of the receiver. THE PICTURE INFORMATION IS RECOVERED FROM THE PICTURE CARRIER BY AN AM DEMODULATOR, AND THE SOUND INFORMATION IS RECOVERED FROM THE SOUND CARRIER BY AN FM DEMODULATOR.

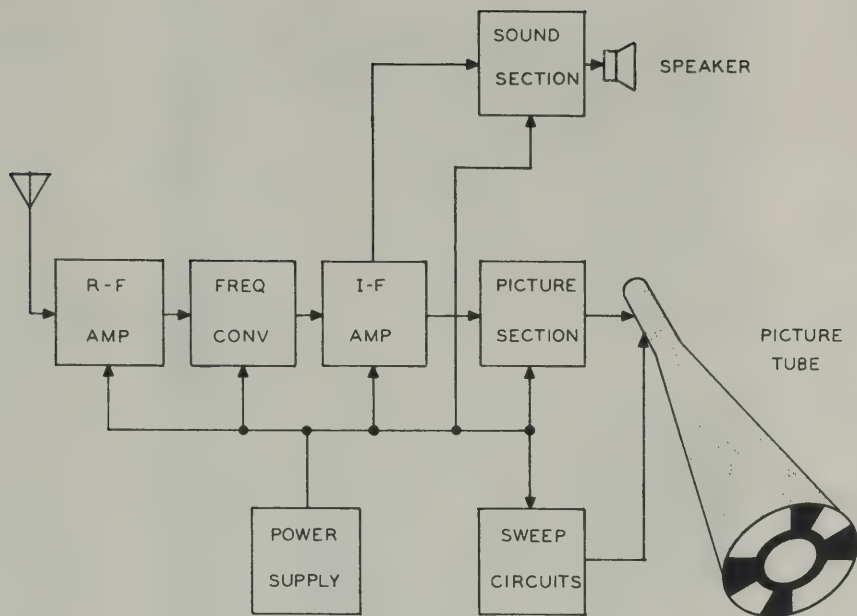


Figure 15

A typical sound section contains one or more additional i-f amplifiers, an FM demodulator, and voltage and power amplifiers just like an ordinary FM radio receiver. The picture section contains an AM demodulator and one or more picture or video amplifiers. The sweep circuits are made up of oscillators and amplifiers which provide the currents required to sweep the electron beam back and forth as well as up and down on the face of the picture tube, thus forming a pattern of lines on which the picture is developed.



A radar receiver section and indicator are combined in this compact unit designed for maritime use. The display section uses a cathode ray tube, or "scope", that is called a plan-position indicator (PPI).

Courtesy Raytheon Company

RADAR SYSTEMS

The term **RADAR** is a contraction of the terms "RADio Detecting And Ranging". Thus, a radar system uses radio waves to determine the direction and range or distance of an object. Basically, the operation of a radar system depends on three simple characteristics of radio waves. These are: (1) radio waves can be focused into narrow beams of energy by means of directive antennas; (2) radio waves travel at a known speed of 300,000,000 meters or 186,000 miles per second; and (3) radio waves are reflected by conductive objects.

Thus, by measuring the time required for a radio wave to travel from a transmitter to an object and be reflected back to a receiver, the distance traveled by the radio wave can be calculated. The direction of the object can be determined by noting the direction in which the narrow beam of radio energy is transmitted and received.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

COMMUNICATION SYSTEMS

RECEIVERS

TELEVISION SYSTEMS

22. The ability of a receiver to select one particular frequency and to reject all others is called (a) selectivity, (b) sensitivity.
23. The ability of a receiver to adequately amplify small signals is called (a) selectivity, (b) sensitivity.
24. The selectivity of a radio receiver is provided by what type of circuits?
25. What name is given to a circuit which recovers the modulating signal from the modulated r-f carrier?
26. Most home entertainment receivers in use today are (a) TRF receivers, (b) superheterodyne receivers.
27. In what type of receiver is the carrier frequency converted to a fixed intermediate frequency before being demodulated?
28. The most common intermediate frequency used in superheterodyne receivers designed to operate in the commercial AM broadcast band is (a) 455 kHz, (b) 10.7 MHz.
29. Why is the wiring and component layout in the i-f amplifiers of an FM receiver more critical than the wiring and component layout in the i-f amplifiers of an AM receiver?
30. What two functions do cathode ray tubes perform in a complete television system?
31. What determines the amount of light emitted by the fluorescent screen on a picture tube?
32. In commercial television broadcasting the sound carrier is (a) amplitude modulated, (b) frequency modulated.
33. In commercial television broadcasting the picture carrier is (a) amplitude modulated, (b) frequency modulated.
34. All television receivers in use today are (a) TRF receivers, (b) superheterodyne receivers.
35. What is the function of the sweep circuits in a television receiver?

Figure 16 illustrates the operation of a simple radar system. The transmitter periodically transmits short bursts or pulses of microwave (very short wave or high frequency) radio energy. These pulses of energy are concentrated into a narrow beam by a directional antenna often called a "dish" because of its shape. Objects called targets in the path of the radio waves reflect a small portion of the energy back to the antenna. A special switch, called a duplexer, permits the same antenna dish to both transmit and receive the r-f energy. The antenna is connected to the transmitter section of the radar system just long enough to transmit a short pulse, or burst, of energy. The antenna is then connected to the receiver section of the radar system by the duplexer and "listens" for the return echo.

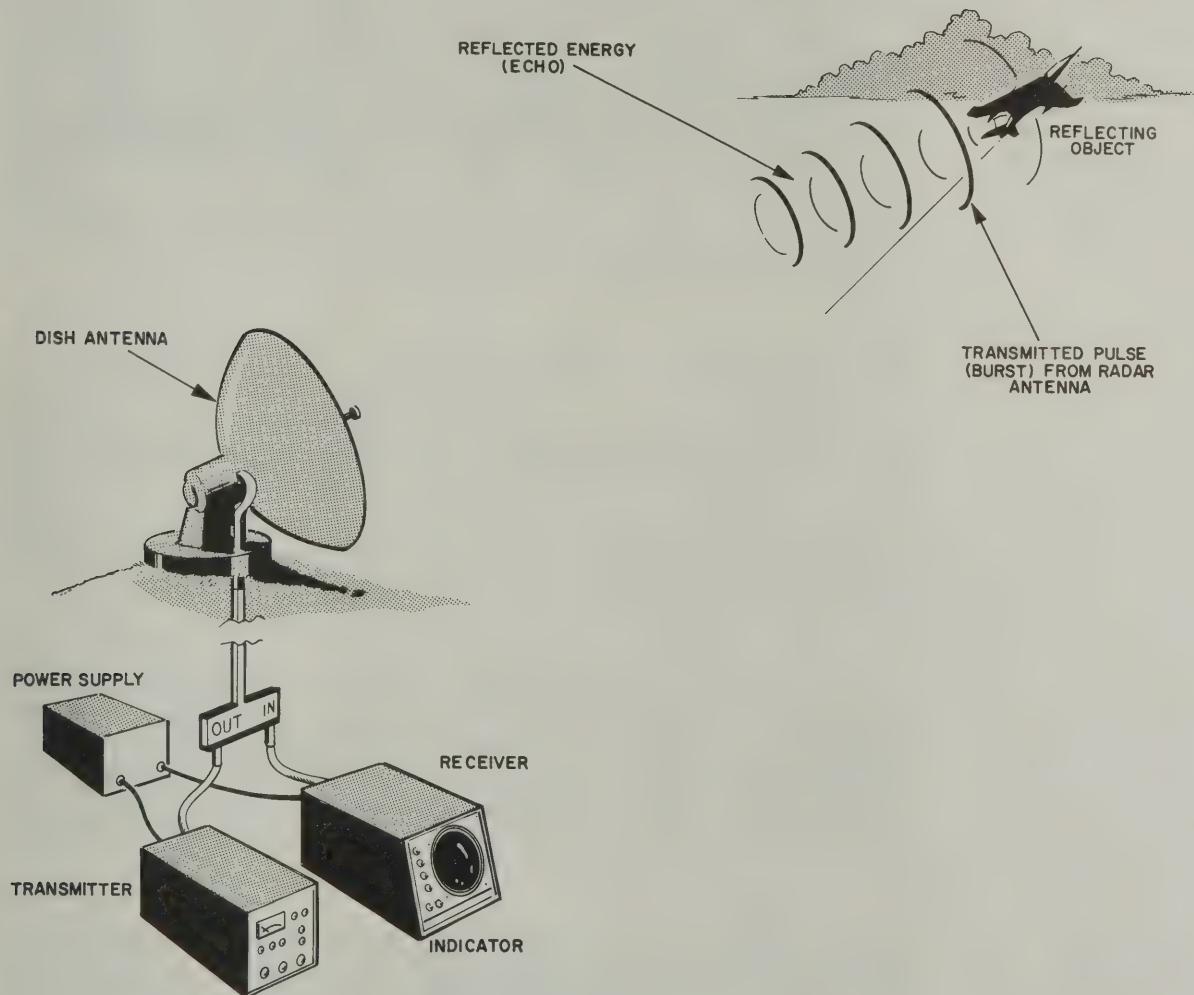


Figure 16

There is a comparatively long time between pulses of transmitted energy to enable the receiver to pick up the reflected waves. If the radio waves were

transmitted continuously, the transmitted waves would blank out the relatively weak reflected waves. Also, since the transmitted pulses have a relatively short duration (about one microsecond), the peak power may be very great without exceeding the power ratings of the transmitter circuits and components. The peak power of a radar transmitter is commonly hundreds of kilowatts.

The distance between the radar station and the reflecting object is determined from a pattern displayed on a cathode ray tube. This cathode ray tube is similar to the picture tube in a television receiver.

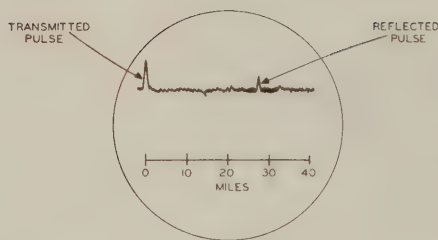


Figure
17

Figure 17 shows one form of display. The electron beam is moved across the face of the cathode ray tube at a constant speed. When a pulse of energy is transmitted by the transmitter, a small pulse of current is applied to the deflection plates of the cathode ray tube. This pulse of current momentarily deflects the electron beam upward as shown in Figure 17. At the end of this pulse, the electron beam moves downward again and continues to move to the right at a constant speed. When the reflected pulse is received, the electron beam is again momentarily deflected in a vertical direction by a voltage pulse applied to the deflection plates.

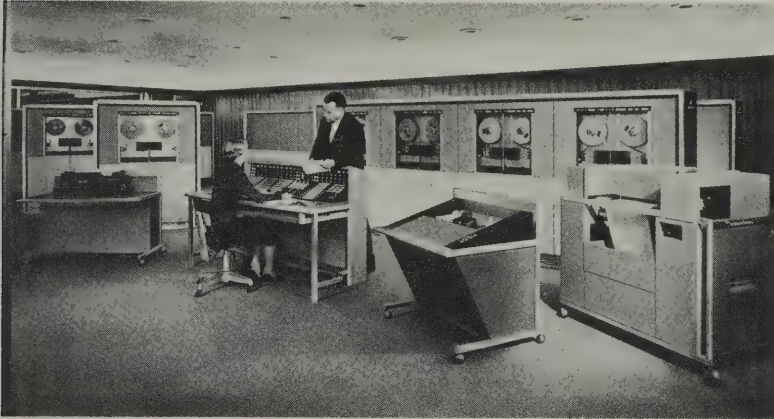
Since the electron beam moves from left to right at a constant speed, the distance between the transmitted and the reflected pulses is a measure of the time required for a pulse to be transmitted to the reflecting object and be reflected back to the receiver. Since radio waves travel at a constant speed in any given medium, the time duration between the transmitted and received pulses is, in turn, a measure of the distance between the radar set and the reflecting object. Thus, scales on the radar indicator may be calibrated in units of distance as shown in Figure 17.

The time duration between the transmitted and the reflected pulse indicated in Figure 17 is the time required for the radio wave to travel from the transmitter to the reflecting object and back to the receiver. This is TWICE the distance from the transmitter to the target. Thus, the distance found by multiplying the speed of the radio wave by the time duration must be divided by 2 to obtain the distance from the transmitter to the target.

For example, assume that the radio wave travels from the transmitter, to the reflecting object, and back to the receiver, in 300 μsec . The speed of the radio wave is 186,000 miles per second. Therefore, the total distance it travels to the reflecting object and back is found by multiplying 186,000 by the time duration:

$$186,000 \times 300 \times 10^{-6} = 55.8 \text{ miles.}$$

The distance from the radar set to the reflecting object is half this distance, or $55.8/2 = 27.9$ miles. The scale on the radar screen is calibrated in terms of distance between the target and the antenna.



Computer systems account for a large percentage of the digital and hybrid systems currently in use. Modern businesses rely upon data processing machines such as this one to record and maintain a variety of business transactions.

Courtesy Radio Corporation of America

INDUSTRIAL ELECTRONIC SYSTEMS

Almost all types of electronic equipment are used in industry. Radio is used for communications; television is used in plant protection; telemetering is used to measure quantities at a distance; and computers are used in research labs and in offices to make out payrolls, for billing and inventory, for calculations and in many other data processing activities. However, the term industrial electronics usually applies to equipment used to control or monitor manufacturing processes, to certain high power equipment used directly to perform some operation on a product being manufactured, and to a few other types of equipment which are very useful in supporting the manufacturing process, but are not directly related to it.

The equipment used to control or monitor manufacturing processes includes electronic timers, temperature controls, motor and generator controls, photo-electric controls, electronic counters, automatic lighting controls, and x-ray equipment. High power equipment used directly to perform some operation on the manufactured product includes r-f heaters, resistance welders, and ultrasonic equipment. This equipment may also require a wide variety of electronic controls. Supporting equipment includes burglar alarms and electronic air cleaners.

Like other types of electronic systems, industrial electronic systems require a wide variety of power supplies, amplifiers, and oscillators. Figure 18 illustrates a typical application of an amplifier in an industrial control circuit. Here an amplifier is used to operate an electromagnetic relay, which in turn, opens and closes a high power circuit.

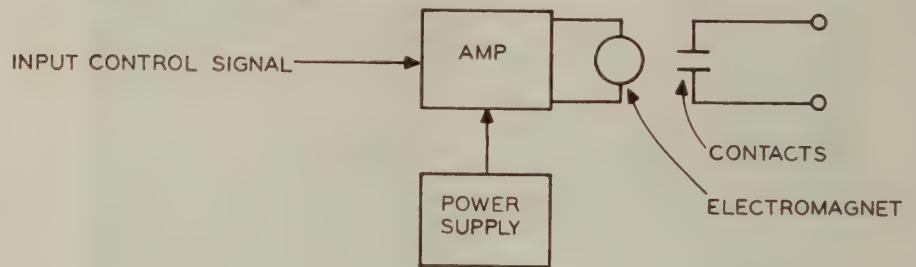


Figure 18

The electromagnetic relay is made up of an electromagnet and a set of contacts. The circle connected to the amplifier output represents the electromagnet and the short horizontal lines to the right of the electromagnet represent the contacts. One of the contacts is stationary and the other is attached to a movable armature. When a current is passed through the electromagnet, the resulting magnetic field attracts the movable armature. As the armature moves, the contact mounted on it moves and closes the circuit.

In applications where the input control signal is capable of supplying a large current, it is applied directly to the electromagnet and no amplifier is required. However, if the input control signal is too weak to operate the relay directly, it must be amplified before being applied to the relay, as shown in Figure 18.

INFORMATION SYSTEMS

Information systems are generally those which facilitate communication between Man and Machine, and thus account for a large percentage of the **DIGITAL SYSTEMS** now in use. The term **MACHINE** is used in reference to data and computation machines. Information systems are among the largest and most complex systems in existence. Most of these are almost totally digital in nature, such as the large and powerful computers. Others, such as spacecraft and flight control systems, are so large that giant computers are merely subsystems of a larger **HYBRID SYSTEM**. Space and aerospace systems require the use of communication, radar, and information systems integrated into a large complex network where there must be man-

to-man, man-to-machine, and machine-to-machine communications. In systems such as this, even machines communicate with each other. Their language, however, is the only language that computers understand—the language of digits.

Computer Systems

COMPUTERS are used to operate upon intelligent information in **DIGITAL** form, and therefore use digital circuitry. In this section we will interpret “Computer” to mean “Digital Computer”. Recall that the communication systems discussed previously were used to transmit intelligent information in **ANALOG** form. This was because we were dealing primarily with man-to-man communications. Computers, however, operate internally upon digits, and thus all communications with machines must ultimately be converted to digital form.

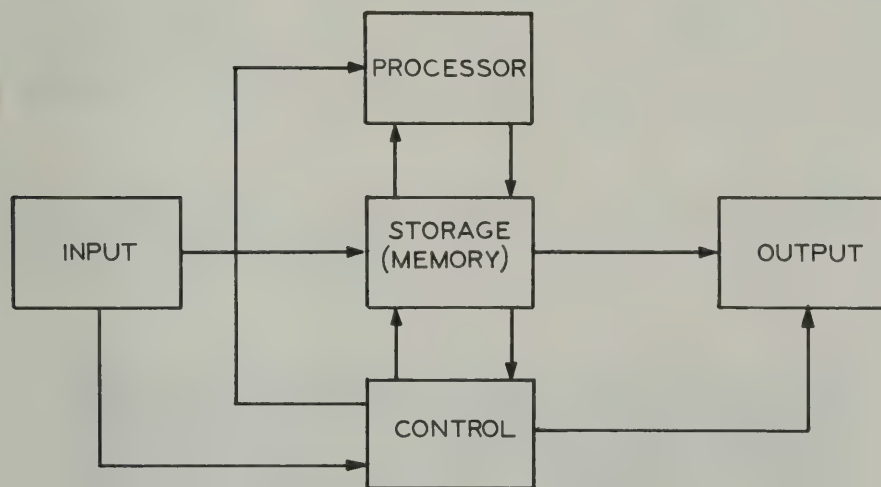


Figure 19

Figure 19 is a block diagram of a typical digital computer system. Each block represents a section, or actually a subsystem, of one complete computer. Although the specific design of each subsystem may vary from one manufacturer to another, all computers follow this basic structure. Note the function of the five basic sections:

1. **INPUT** — the section that provides the method for placing information into the computer.

2. **STORAGE** — the section which retains the information received by the computer until it is required; also referred to as the “memory”.
3. **CONTROL** — the section which sequences and supervises the computer operations, telling the computer what to do and when to do it.
4. **PROCESSOR** — the section which performs the arithmetic operations upon the information placed in the computer.
5. **OUTPUT** — the section that provides a means for obtaining the results, or answers, from the computer and making them available to the user.

The actual electronic circuits which make up the sections, and the overall system, are known as computer **HARDWARE**. These are the elements which perform logic functions, route signals, and also include mechanical devices such as an output printer. The five sections of Figure 19 cannot function, however, without **SOFTWARE**, called a **COMPUTER PROGRAM**. The program is the link between man and computer. Since a computer cannot “think” by itself, it must be told what to do, and in what order to do it. The computer program does this. It is nothing more than a list of operations telling the computer exactly what to do.

Computers are classified as to the type of operation which will be required during use. That is, the system is either **GENERAL PURPOSE**, or **SPE-**



Digital and hybrid systems provide communication links between man and machine. This system reads ordinary typewritten or printed information and translates it into computer language. The Computing Reader includes a general-purpose computer as a subsystem.

Courtesy Recognition Equipment, Inc.

GENERAL PURPOSE. A computer with a large memory and multiple programming capabilities can be used for a variety of purposes, and is a general purpose computer. A computer which performs only a limited number of functions, or perhaps only one function but one which is performed repetitively, is a special purpose computer. Special purpose computers generally operate at higher speeds than general purpose systems. This is because the hardware can be designed to execute the instruction in the program in a very efficient manner. Since the designers know that the computer will be used to solve only one type of problem, much of the circuitry required for a general purpose computer can be eliminated.

Information Storage and Retrieval Systems

Information storage and retrieval systems vary considerably in type, size and operation. They may be as simple as a single filing cabinet, or as complex as a system of computers. Some storage and retrieval systems use photographic film (microfilm) to store information. Others, such as the one discussed in this lesson, use television recording equipment called a **VIDEO TAPE RECORDER** or **VTR**. The system shown in Figure 20 is typical of large storage and retrieval systems, referred to as simply retrieval systems or **FILES**.

Retrieval systems are organized in one of two ways: they are either **SERIAL ACCESS** or **RANDOM ACCESS**. This refers to the method used for **ENTERING** the file. Thus, the technique used for gaining entry to a serial access file is called **INDEXING**. The technique used for gaining entry to a random access file is called **RANDOMIZING**. A simple example which illustrates indexing is the way in which an ordinary filing cabinet is organized. In order to recover a document from a filing cabinet, the document must be placed in a specific location so that it can be found again. Usually, this is according to some system, such as keeping the file in alphabetical or numerical order. In other words, documents are filed in some type of serial order. In this way, a document titled "transistors" may be found by going to the "T's", and looking only through the documents beginning with T. Suppose, however, the filing clerk had paid no attention to the manner in which the document was filed, and had placed it just anywhere in the cabinet. This would mean that in trying to find the document again, all the documents, beginning with A, would have to be searched until the desired one was found. Such a system would be called a "random file," and although placing information **IN** the file would be quite easy, retrieval would be inconvenient. This is especially true if a large file requiring several filing cabinets is used.

Searching for a document by using an electronic system such as shown in Figure 20 is done in the same way as would be required by a file clerk. That is, since the document could be anywhere in the file, each document must

be "looked at" starting at the front of the file. The difference, however, is that the electronic system is capable of searching at electronic speeds. The electronic system thus permits the convenience of random filing while retaining the ability to retrieve the document easily.

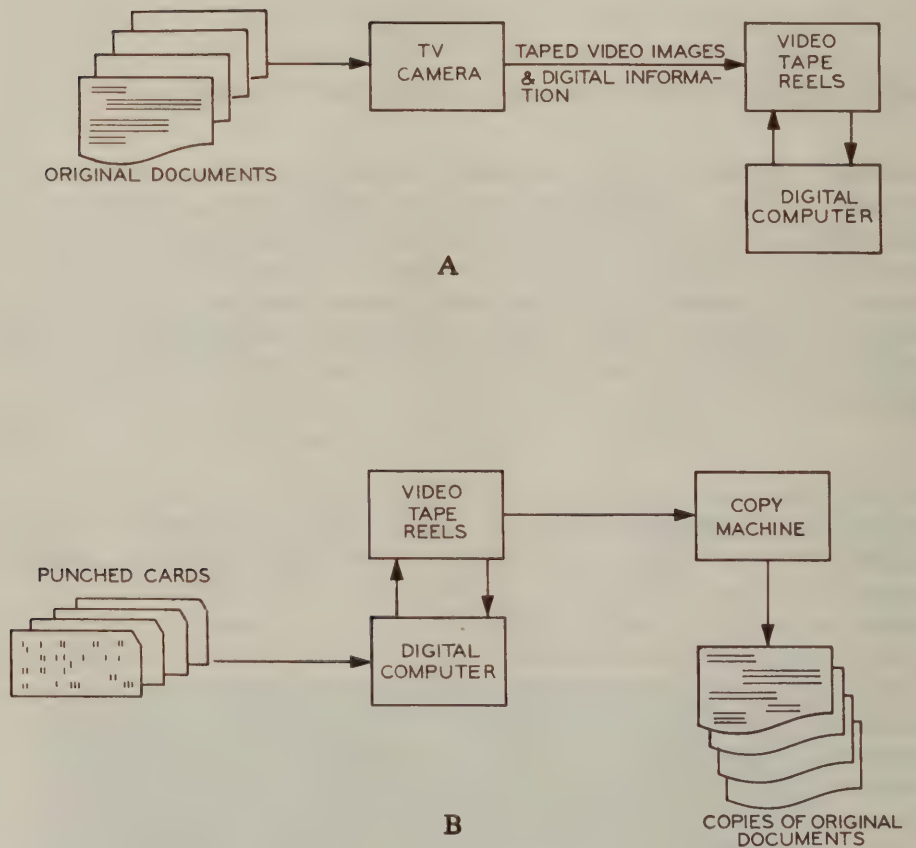


Figure 20

Systems such as the one in Figure 20 might be used by institutions and merchandising firms where records of daily transactions must be kept and filed. Consider, as a hypothetical example, this system being used in a school or university. Large universities must maintain extensive files of information concerning the students who attend. Among these are financial records, grades, personal data, campus addresses, and a variety of other information. In order for systems such as this to function, each student must be assigned a student number. This may be the number on his ID card, or some other

The following Practice Exercise questions cover the subjects which you have just studied. They are:

RADAR SYSTEMS

INDUSTRIAL ELECTRONIC SYSTEMS

INFORMATION SYSTEMS

COMPUTER SYSTEMS

INFORMATION STORAGE AND RETRIEVAL SYSTEMS

36. A radar transmitter transmits radio waves (a) continuously, (b) in short pulses.
37. If a radar pulse requires one millisecond to be transmitted to a target and reflected back to a receiver, the distance between the target and the radar set is (a) 93 miles, (b) 186 miles.
38. Why is it sometimes necessary to precede a relay with an electronic amplifier in an industrial control circuit?
39. The control section of a digital computer performs the arithmetic operations upon the input information. True or False?
40. A list of instructions which describes the operations to be performed by a computer is called a _____ .
41. A special purpose computer can execute instructions faster than a general purpose computer. True or False?

number which can be used by the electronic system. A record concerning the student, such as a tuition receipt, is filed in the electronic system as shown in Figure 20A. In this case, the receipt is represented by one of the "original documents" at the left. An operator photographs the document on video tape. (This procedure is similar to the video tape "instant replay" technique used in commercial TV broadcasting.) In addition to the video image, however, the student's ID number and other coded information may be placed on the tape in digital form, and stored on the video tape reels. The document is now permanently filed in the system, and the operator may place another document in the system by simply advancing the video tape. Note that the document did not have to be placed on any specific portion of the tape. The operator merely placed it on the first blank segment found by advancing the tape. Thus, the electronic system is a random file, and is also referred to as **OPEN ENDED**. This means that information can be added to either end of the file, front or back. Actually, additions are not limited to just the ends. Old information can be erased from any portion of the tape, and new information inserted in its place.

The information placed in the file can be retrieved as shown in Figure 20B. Using the example above, assume that the student requests a duplicate copy of the tuition receipt filed earlier. The operator asks the student for his student number, and punches the number and a code number corresponding to "financial information" on a data card. This is represented by one of the punched cards shown at the left in Figure 20B. The card is fed into the computer which starts searching the tape reels by advancing them at high speed. The computer must continue to search the video tape until the digital information placed on the tape matches the information on the punched card. As pointed out earlier, this is equivalent to searching the entire file, beginning at the front. The computer, however, does this at very high speed. Thousands of documents per second can be searched by electronic systems such as this.

When the document is found, the tape stops advancing. A video tape "picture" of the document is then converted to a paper copy by the copy machine, and the copy is given to the student. The average time elapsed during a search involving hundreds of thousands of documents is typically less than 10 seconds for electronic information retrieval systems.

SUMMARY

Electronic systems are made up of a few basic types of circuits connected together. Three of the most common basic circuits are amplifiers, oscillators, and switches. Amplifiers and oscillators are used to provide analog functions where the circuit output can be varied smoothly over a continuous range. Analog circuits use active elements (electron tubes or transistors). Switches provide a digital function where the output may be only on or off. Digital circuits may use an active element, but active elements are not always required. A simple digital circuit may use switches or relays and other passive elements. Amplifiers are used to increase the magnitudes of electric signals. The ratio between the output and the input of an amplifier is called the gain of the amplifier. Oscillators are used to generate alternating currents and voltages in electronic systems. An oscillator is made up of an amplifier and a feedback circuit which feeds a portion of the output back to the input.

The main types of radio communication systems are radiotelegraphy, used to transmit coded messages; radiotelephony, used to transmit voice, music, and other sounds; television, used to transmit moving pictures; and facsimile, used to transmit still pictures. In all of these systems, information is transmitted by impressing it on an r-f carrier wave.

The process of impressing information on an r-f carrier wave is called modulation. The two main types of modulation are amplitude modulation, in which the amplitude of the r-f carrier is varied in accordance with the modulating signal, and frequency modulation, in which the frequency of the r-f carrier is varied in accordance with the modulating signal. Continuous wave radiotelegraphy is actually a form of amplitude modulation since the amplitude of the carrier varies from zero to some fixed amplitude in forming a dot or a dash.

In television broadcasting the sound carrier is frequency modulated and the picture carrier is amplitude modulated. Cathode ray tubes are used to convert the picture to an electric signal and to convert the electric signal back into a picture. As the electron beam in the camera tube sweeps back and forth across the picture, the electron beam in the picture tube sweeps back and forth across a fluorescent screen and "paints" the picture on the screen.

Television sets are made up of superheterodyne receivers with some provision for separating the sound and picture signals and with sweep circuits which provide the currents required to sweep the electron beam back and forth across the fluorescent screen.

In a radar system, pulses of high frequency radio energy are transmitted by a transmitter. Portions of these pulses are reflected back to a receiver by objects which are then called targets. The time required for a transmitted pulse to reach an object and be reflected back to the receiver is a measure of the distance between the radar set and the reflecting object.

Industrial electronic systems vary widely in type and application. Manufacturing industries require a variety of monitoring and central equipment. These are often associated with automated production devices such as timers, temperature controls, counters, and other electronic aids.

Information systems often use digital circuitry, and account for a large number of the digital systems in use. However, since information systems are designed to aid man-to-man, man-to-machine, and machine-to-machine communications, many are hybrid systems.

Most of the computers in use are digital systems, and consist of five basic sections: input, storage, control, processing, and output. Computers are used with the aid of a computer program, which determines the computer operating sequences. Computers are either general purpose or special purpose. A general purpose computer is flexible, and may be used with several types of programs. A special purpose computer is relatively inflexible, and often is built to accept only one type of program, but it usually can function at a higher rate of speed than can a general purpose computer.

Information storage and retrieval systems are either serial access or random access. Ordinary filing cabinets are an example of serial access storage and retrieval systems which require indexing as the mode of entry. Electronic storage and retrieval systems are capable of random access operation, using random modes of entry.

IMPORTANT DEFINITIONS

ACTIVE ELEMENT — An electronic component capable of providing amplification — such as an electron tube or a transistor.

AMPLIFICATION — One of the three basic circuit functions made possible through the use of an electron tube or transistor.

AMPLIFICATION FACTOR — See GAIN.

AMPLITUDE MODULATION (AM) — A type of modulation in which the amplitude of the carrier varies in accordance with the modulating signal.

ANALOG CIRCUIT — A circuit using an active element whose output is proportional to the input.

ANALOG SYSTEM — A number of analog circuits used together to operate upon information in analog form.

BANDWIDTH — The range of frequencies over which a circuit provides relatively constant output with a constant input. Thus, bandwidth is the range of frequencies over which an electronic amplifier has relatively constant gain. ($\pm 3\text{dB}$.)

BISTABLE ELEMENT — An active or passive element connected in a circuit which is capable of producing one of only two output states, i.e. "on" or "off".

CLOSED LOOP — A circuit using feedback techniques to provide control of the output signal.

COMPUTER — A system used to operate upon information which is capable of performing simple operations or calculations at high speed.

COMPUTER PROGRAM — A detailed list of instructions which describes the type and the order of operations to be performed by a computer.

DEMODULATOR — A circuit used to recover the modulating signal from a modulated carrier. Also called a DETECTOR.

DETECTOR — See DEMODULATOR.

DIGITAL CIRCUIT — A circuit which may use an active or passive element whose output switches states or levels in response to the input signal.

DIGITAL SYSTEM — A number of digital circuits used together to operate upon information in digital form.

IMPORTANT DEFINITIONS (Continued)

DISTORTION — The differences between the input and output waveforms of an amplifier other than an increase in amplitude and an inversion of polarity.

FEEDBACK — An analog circuit function in which a portion of a signal voltage is returned to a preceding point in the circuit.

FILE — A collection of stored information arranged for retrieval of any single item at some later time.

FLIP-FLOP — A bistable element used extensively in digital circuits.

FREQUENCY MODULATION (FM) — A type of modulation in which the frequency of the carrier varies in accordance with the modulating signal.

GAIN — The increase in amplitude provided by an amplifier. The ratio between the output and the input. Also called the **AMPLIFICATION FACTOR**.

HARDWARE — Often referred to as the physical circuitry and mechanical structure of a computer system.

HIGH FIDELITY — A high degree of faithfulness in reproducing an input signal. A low degree of distortion.

HYBRID SYSTEM — A number of analog and digital circuits used together to operate upon information which may be in either analog or digital form.

INDEXING — The method used for gaining entry to a serial access file.

INTERMEDIATE FREQUENCY (i-f) — The frequency to which the received signal is converted in a superheterodyne receiver.

MODULATION — The process of impressing information on a carrier.

OPEN ENDED — A random access file in which new information may be inserted anywhere in the file.

OSCILLATION — A basic circuit function made possible through the use of an active element and a continuous feedback loop.

OSCILLATOR — A circuit used to generate alternating voltages and currents.

PASSIVE ELEMENT — An electronic component incapable of providing amplification, such as a resistor, a capacitor or an inductor.

IMPORTANT DEFINITIONS (Continued)

RADAR — A contraction of the terms RAdio Detecting And Ranging. The use of radio to determine the direction and range, or distance, of an object.

RADIOFACSIMILE — The transmission and reception of printed or photographic material via radio waves with a copy being obtained at the receiver.

RADIOTELEGRAPHY — A type of radio communication system used to transmit coded messages in the form of dots and dashes.

RADIOTELEPHONY — A type of radio communication system used to transmit voice, music and other sounds.

RANDOM ACCESS — A type of information storage and retrieval system where no specific file order is required.

RANDOMIZING — The method used for gaining entry to a random access file.

SELECTIVITY — The ability of a radio receiver to select one particular frequency and to reject all others.

SENSITIVITY — The ability of a radio receiver to amplify small signals.

SERIAL ACCESS — A type of information storage and retrieval system which requires that each item in the file be placed in a specific place.

SIGNAL — Voltages or currents which vary in accordance with some form of information.

SOFTWARE — The means and techniques used to program a computer.

STAGE — An active element and its associated passive elements which perform some function in an electronic system.

SUPERHETERODYNE RECEIVER — A type of radio receiver in which all r-f carriers selected by the receiver are converted to the same intermediate frequency before demodulation.

SWITCHING — One of the basic circuit functions and one which provides one of only two possible output states.

TELEVISION — The continuous transmission and reception of images via radio waves which when viewed provides continuous visual communication.

IMPORTANT DEFINITIONS (Continued)

TUNED RADIO FREQUENCY (TRF) RECEIVER — A type of radio receiver in which radio frequency amplifiers must be retuned each time a different r-f carrier is selected.

VIDEO TAPE RECORDER (VTR) — A magnetic tape recording system capable of recording and storing television signals.

PRACTICE EXERCISE SOLUTIONS

1. active
2. True — A simple switch or relay will function in a switching circuit.
3. (a) oscillator. An oscillator circuit is closed loop since feedback is used.
4. (c) bistable elements. A bistable element is the equivalent of a simple switch.
5. True — Electron tubes and transistors may be used in circuits which do not amplify or oscillate, but merely transfer a signal from one circuit to another.
6. distortion. — In most amplifiers, such as those used to amplify audio signals, it is desirable to keep distortion to a minimum. In some industrial amplifiers, however, the main concern is to develop enough power to operate some output device such as a motor or solenoid. In these amplifiers, the important thing is the increase in amplitude and the distortion is of no importance.
7. A high fidelity amplifier is capable of producing an amplified version of the input signal with a high degree of faithfulness; that is, with very little distortion.
8. 20 — The gain of an amplifier is the ratio of the output to the input. Thus, $20/1 = 20$.
9. Bandwidth is the range of frequencies over which an amplifier has a relatively constant gain.
10. (a) an audio amplifier. — Video amplifiers have a relatively constant gain from a few hertz to over 4 megahertz.
11. (a) an oscillator — A rectifier converts ac to dc.
12. An amplifier and a feedback circuit.
13. It must have the proper phase and amplitude.
14. False — The output of an oscillator may have a wide variety of waveforms including sawtooth, triangular, and rectangular waves.
15. (1) Coded messages in the form of dots and dashes, (2) voice, music, and other sounds, (3) moving pictures from film or tape or live from a studio, and (4) still pictures transmitted from one location to another. These four types of radio communication systems are called radio-telegraphy, radiotelephony, television, and facsimile, respectively.

PRACTICE EXERCISE SOLUTIONS (Continued)

16. An oscillator.
17. modulation.
18. True — In radio communication systems, information is impressed on an r-f carrier by changing the r-f carrier in accordance with the information to be transmitted.
19. (c) an amplifier circuit — The audio frequency amplifier used in this application is called a modulator.
20. The amplitude of the r-f carrier. — The term AM stands for amplitude modulation.
21. The frequency of the r-f carrier. — The term FM stands for frequency modulation.
22. (a) selectivity.
23. (b) sensitivity.
24. Tuned or resonant circuits.
25. Demodulator or detector — A common circuit used for demodulating AM carriers is much like the rectifier circuits used in power supplies.
26. (b) superheterodyne receivers. — Superheterodyne receivers are much easier to tune than TRF receivers.
27. A superheterodyne receiver.
28. (a) 455 kHz — 10.7 MHz is the most common i-f used in superheterodyne receivers designed to operate in the commercial FM broadcast band.
29. Because radio frequency effects are much more pronounced at the intermediate frequency of 10.7 MHz used in FM receivers than at the intermediate frequency of 455 kHz used in AM receivers.
30. A cathode ray tube called a camera tube converts the picture to be televised to an electric signal. Another type of cathode ray tube called a picture tube converts the electric signal back to a picture.
31. The intensity of the electron beam, which is controlled by the electric signal originally developed by the camera tube.
32. (b) frequency modulated.

PRACTICE EXERCISE SOLUTIONS (Continued)

- 33. (a) amplitude modulated.
- 34. (b) superheterodyne receivers.
- 35. To provide the currents required to sweep the electron beam back and forth as well as up and down on the face of the picture tube.
- 36. (b) in short pulses. — The radar receiver picks up the reflected waves during the time between transmissions. Also, transmitting short pulses allows the peak power to be much greater than would be possible if the transmitter transmitted continuously.
- 37. (a) 93 miles — Since electromagnetic waves travel at a speed of about 186,000 miles per second, a radar pulse travels $186,000 \times .001$, or 186 miles in one millisecond. However, this is the total distance the radar pulse travels, from the radar set to the target and back again. The distance from the radar set to the target is one-half of 186 miles, or 93 miles.
- 38. Because the control signal may not have sufficient amplitude to operate the relay and therefore must be amplified.
- 39. False — The Processor section performs the arithmetic operations, not the control section.
- 40. computer program.
- 41. True — A special purpose computer can be designed with special circuits which perform only one type of operation, and may be designed to perform in the fastest possible time.

De VRY INSTITUTE OF TECHNOLOGY
4141 BELMONT AVENUE, CHICAGO, ILLINOIS 60641

ONE OF THE
BELL & HOWELL SCHOOLS
I085A

I085A
EXAMINATION
CHECK SHEET

1. B - ONE CHARACTERISTIC OF AN ANALOG CIRCUIT IS THE USE OF -- FEEDBACK.

Feedback is the technique by which the output of a circuit is controlled by the input, which is a true analog function.

2. B - IF A CERTAIN AMPLIFIER DEVELOPS A SIGNAL OUTPUT OF 30 VOLTS WITH A SIGNAL INPUT OF 1 VOLT, THE GAIN OF THE AMPLIFIER IS -- 30.

Amplifier gain is defined as the ratio of the output signal to the input signal. Thus $30/1 = 30$.

3. B - THE R-F CARRIER IN A RADIO TRANSMITTER IS USUALLY GENERATED BY -- AN ELECTRONIC OSCILLATOR.

An r-f carrier is the high frequency alternating voltage waveform which is produced by an active element capable of producing and sustaining high frequency oscillations.

4. D - THE TRANSMISSION OF CODED MESSAGES IN THE FORM OF DOTS AND DASHES BY MEANS OF RADIO WAVES IS CALLED -- RADIO TELEGRAPHY.

The system of dots and dashes as a code, originally over wires, is called telegraphy. Therefore, the use of the same system using radio waves is called radiotelegraphy.

5. C - THE PROCESS OF TRANSMITTING INFORMATION BY VARYING THE AMPLITUDE OF AN R-F CARRIER WAVE IS CALLED -- AMPLITUDE MODULATION.

The r-f carrier, which is fixed in frequency and amplitude, has impressed upon it a VARYING AMPLITUDE signal, which is the intelligent information recovered at the receiver.

6. B - THE ABILITY OF A RADIO RECEIVER TO SELECT ONE PARTICULAR FREQUENCY AND REJECT ALL OTHERS IS CALLED -- SELECTIVITY.

A radio receiver must be able to select one particular frequency at a time out of many which are transmitted in a complex radio network.

7. C - THE TYPE OF RADIO RECEIVER IN WHICH THE INCOMING R-F CARRIER IS CONVERTED TO AN INTERMEDIATE FREQUENCY BEFORE DEMODULATION IS CALLED -- A SUPERHETERODYNE RECEIVER.

A superheterodyne receiver converts the incoming r-f carrier frequency to an intermediate frequency so that no matter what r-f carrier frequency is received, the intermediate frequency is always the same.

8. B - IN A TELEVISION RECEIVER, THE PICTURE INFORMATION IS RECOVERED FROM THE PICTURE CARRIER BY -- AN AM DEMODULATOR.

The picture, or video, portion of television transmission is amplitude modulated and is demodulated separately from the sound portion of the signal using an AM demodulator.

9. B - IF A RADAR PULSE TRAVELS FROM THE TRANSMITTER TO A REFLECTING OBJECT IN 200 μ SEC, THE DISTANCE FROM THE RADAR SET TO THE OBJECT IS -- 37.2 MILES.

Since the pulse travels to the object at 186,000 miles per second and requires 200 μ sec in one direction, the pulse travels the round trip from antenna - object - antenna in 400 μ sec. Thus $186,000 \times 400 \times 10^{-6} = 74.4$ miles (the round trip distance). However, the object itself is $1/2$ this distance: $74.4/2 = 37.2$ mi.

10. B - THE INFORMATION RETRIEVAL SYSTEM SHOWN IN FIGURE 20 IS A -- HYBRID SYSTEM.

The information system uses both analog and digital recording techniques, and requires SUBSYSTEMS of both analog and digital circuits.

PRACTICE EXERCISE SOLUTIONS (Continued)

- 33. (a) amplitude modulated.
- 34. (b) superheterodyne receivers.
- 35. To provide the currents required to sweep the electron beam back and forth as well as up and down on the face of the picture tube.
- 36. (b) in short pulses. — The radar receiver picks up the reflected waves during the time between transmissions. Also, transmitting short pulses allows the peak power to be much greater than would be possible if the transmitter transmitted continuously.
- 37. (a) 93 miles — Since electromagnetic waves travel at a speed of about 186,000 miles per second, a radar pulse travels $186,000 \times .001$, or 186 miles in one millisecond. However, this is the total distance the radar pulse travels, from the radar set to the target and back again. The distance from the radar set to the target is one-half of 186 miles, or 93 miles.
- 38. Because the control signal may not have sufficient amplitude to operate the relay and therefore must be amplified.
- 39. False — The Processor section performs the arithmetic operations, not the control section.
- 40. computer program.
- 41. True — A special purpose computer can be designed with special circuits which perform only one type of operation, and may be designed to perform in the fastest possible time.

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1085A

Example: Apples grow on

A	
B	
C	
D	

A. vines. B. trees. C. bushes. D. grass.

1. A ☐
 B ☒
 C ☐
 D ☐

One characteristic of an analog circuit is the use of
 (A) passive components. (B) feedback. (C) switching. (D) bistable elements.
2. A ☐
 B ☒
 C ☐
 D ☐

If a certain amplifier develops a signal output of 30 volts with a signal input of 1 volt, the gain of the amplifier is
 (A) 60. (B) 30. (C) 15. (D) 1.
3. A ☐
 B ☒
 C ☐
 D ☐

The R-F CARRIER in a radio transmitter is usually GENERATED BY
 (A) a microphone. (B) an electronic oscillator. (C) a rectifier. (D) a modulator.
4. A ☐
 B ☐
 C ☐
 D ☒

The transmission of coded messages in the form of dots and dashes by means of radio waves is called
 (A) radio facsimile. (B) radiotelephony. (C) television. (D) radiotelegraphy.
5. A ☐
 B ☐
 C ☒
 D ☐

The process of transmitting information by varying the AMPLITUDE of an r-f carrier wave is called
 (A) frequency modulation. (B) amplifying. (C) amplitude modulation. (D) rectifying.
6. A ☐
 B ☒
 C ☐
 D ☐

The ability of a radio receiver to select one particular frequency and reject all others is called
 (A) demodulation. (B) selectivity. (C) fidelity. (D) sensitivity.
7. A ☐
 B ☐
 C ☒
 D ☐

The type of radio receiver in which the incoming r-f carrier is converted to an intermediate frequency before demodulation is called
 (A) a TRF receiver. (B) an oscillator. (C) a superheterodyne receiver. (D) a rectifier.
8. A ☐
 B ☒
 C ☐
 D ☐

In a television receiver, the PICTURE INFORMATION is RECOVERED from the picture carrier by
 (A) an FM demodulator. (B) an AM demodulator. (C) an i-f amplifier. (D) a video amplifier.
9. A ☐
 B ☒
 C ☒
 D ☐

If a radar pulse travels from the transmitter to a reflecting object in 200 μ sec, the distance from the radar set to the object is
 (A) 18.6 miles. (B) 37.2 miles. (C) 9.3 miles. (D) 74.4 miles.
10. A ☐
 B ☒
 C ☐
 D ☒

The information retrieval system shown in Figure 20 is a
 (A) communication system. (B) hybrid system. (C) digital system. (D) serial access file.

VACUUM TUBE HIGH FREQUENCY AMPLIFIERS

1088

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





The r-f amplifier shown above is classed as a "Linear Amplifier". A linear amplifier is an r-f amplifier especially designed to introduce as little distortion as possible and is generally operated Class AB or Class B.

Courtesy Hafstrom Technical Products, Inc.

CONTENTS

Wavelengths	Page 4
Frequency Spectrum	Page 5
High-frequency Electronic Equipment	Page 6
Stage-to-stage Transfer of Energy	Page 7
H-F Coupling Methods	Page 8
Capacitive Coupling	Page 9
Inductive Coupling	Page 16
High-frequency Amplifier Tubes	Page 19
H-F Voltage Amplifiers	Page 19
H-F Power Amplifiers	Page 20
Grid Driving Power	Page 24
Push-Pull Amplifiers	Page 24
Frequency Multipliers	Page 24
Neutralization	Page 26
Plate Neutralization	Page 27
Grid Neutralization	Page 28
The Grounded-grid Amplifier	Page 28
Suppression of Parasitic Oscillation	Page 29

LIMITLESS

*There is nothing, I hold, in the way of work,
That a human may not achieve
If he does not falter, or shrink or shirk,
And more than all, if he will believe.
And whatever the height you yearn to climb
Tho' it never was trod by the foot of man,
And no matter how steep—I say you can,
If you will be patient—and use your time.*

VACUUM TUBE HIGH FREQUENCY AMPLIFIERS

High-frequency voltages and currents are used in almost every phase of modern electronics. They are used as carriers of intelligence in radio, television and communications systems. In radar, high-frequency signals allow us to "see" objects at great distances. Industry employs high-frequency energy to seal material such as plastic, and to heat and harden steel and various alloys. High-frequency waves are also used in diathermy machines used by the medical profession. Electronic ovens that employ high-frequency currents can cook an entire meal in a matter of minutes.

In many applications the high-frequency voltages and currents must be amplified before they can be put to use. This is especially true of radio and television receivers. The outputs from the receiving antennas of such devices are very small. Therefore, the signals must be amplified before they can be used to trace a picture or operate a loudspeaker.

Before examining the circuits used in high-frequency amplifiers, a review of the wavelength concept and the high-frequency energy spectrum will be presented.

WAVELENGTHS

All electromagnetic energy is in the form of waves. Heat, light, and radio waves are forms of electromagnetic energy. They differ only in wavelength. The wavelength of a wave is the distance in space that one cycle of energy occupies. Wavelengths range from a small fraction of an inch to many miles.

Electromagnetic energy travels at the speed of light (300,000,000 meters per second). Thus, if a wave is 1 meter long, 300,000,000 of these waves will pass a given point in one second and the frequency will be 300,000,000 Hz or 300 MHz. If the wavelength is increased from 1 to 2 meters, only half as many waves will pass the point in a second. Thus, the frequency decreases to 150,000,000 Hz or 150 MHz.

From the above examples a simple relation can be noted. Wavelength can be found by dividing the velocity of the wave by its frequency. Using the Greek letter λ (lambda) to represent wavelength, a simple formula expressing wavelength in terms of speed and frequency can be shown as:

$$\lambda = \frac{300,000,000}{f} \quad (1)$$

where

f = frequency in hertz

λ = wavelength in meters

From the above expression it may be noted that as the frequency increases, the wavelength decreases. Thus the shortest wavelengths have the highest frequencies. This expression can be simplified by expressing the frequency in MHz.

$$\lambda = \frac{300}{f(\text{MHz})} \quad (1a)$$

To determine frequency when wavelength is known, the expression can be rearranged as shown below.

$$f(\text{MHz}) = \frac{300}{\lambda} \quad (1b)$$

FREQUENCY SPECTRUM

A chart showing the frequency ranges occupied by the various kinds of electromagnetic energy is shown in Figure 1. The kinds of energy with the longest wavelengths are indicated near the left in this chart, and those with the shorter wavelength energies near the right. As shown, a band designated **"RADIO FREQUENCIES" (R-F)** extends from 1×10^4 to 3×10^{12} Hz.

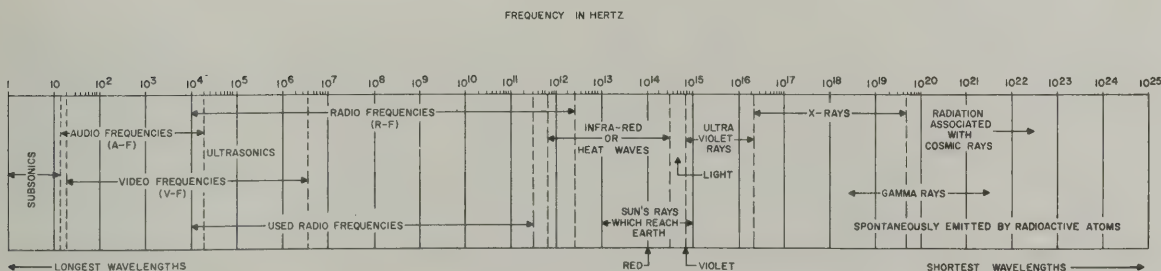


Figure 1

The name Radio Frequencies was given to specify energy which radiates readily through space, and does not apply only to those wavelengths used for

radio broadcast and communication. Actually, the energy with frequencies above the r-f band has this characteristic of easy radiation also. However, the term radio frequencies was adopted many years ago, before much was known of the shorter wavelengths, and is retained today as a means of distinguishing between the various bands. Often the r-f band is called the high-frequency (h-f) band just as the audio band is often called the low-frequency (l-f) band. In this lesson we will use the term h-f to designate those frequencies which are easily radiated through space.

In the chart, each solid vertical line corresponds to a frequency ten times greater than the one to its left. From left to right, the spaces between the lines represent larger and larger frequency bands. Thus, the heat, light, and X-ray bands occupy a much greater portion of the frequency spectrum than the audio and radio frequency waves.

HIGH-FREQUENCY ELECTRONIC EQUIPMENT

A large number of the applications, including high frequencies, are related to communications such as radio, television, radar, two way radio, etc. Also, as mentioned earlier, there are wide fields of industrial, medical, and scientific applications, all of which are as important as the communication applications. In fact, all the circuits used in this frequency range are similar even though their applications may differ. To show this similarity, the general layout of a typical high-frequency control device is shown in Figure 2A. In Figure 2B the general layout of a radio communications system is given for comparison.

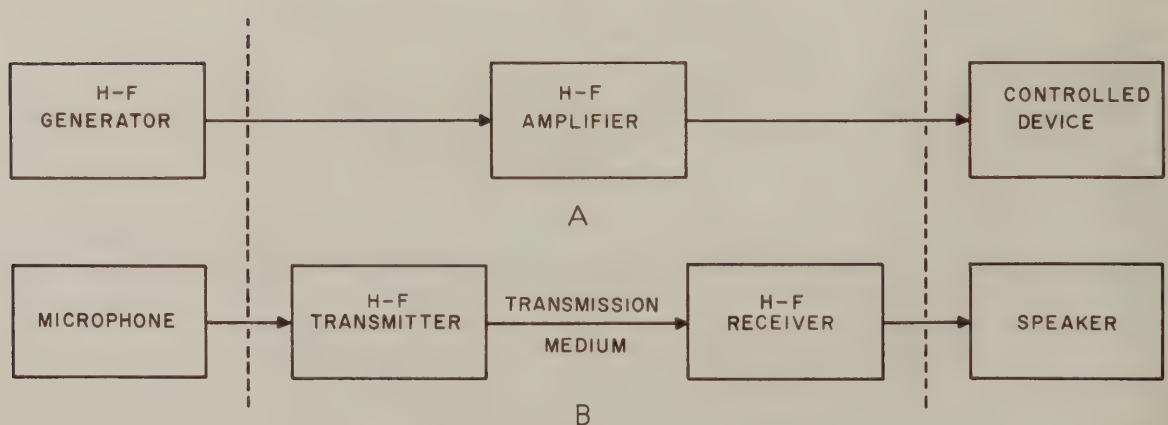


Figure 2

From a broad point of view, both systems may be considered as types of communication systems. In each case, an energy conversion unit produces electric variations which are communicated to another point.

In Figure 2A, the output of an h-f generator is connected directly to the h-f amplifier. In turn, the controlled unit is connected to the amplifier output. Often, the distance between the input and output units is so short that the transmission loss is negligible and only a small amplifying unit is required.

Every communication system is composed of two main parts which may be designated as the TRANSMITTER and the RECEIVER. As shown in Figure 2B, the transmitter is located at the point from which the message or intelligence originates, and the receiver is at the point to which the intelligence is communicated. Between the transmitter and the receiver, the message may be carried by wire, such as a telephone or telegraph line, or radiated through space, as in radio broadcasting. The choice of the type of transmitting medium depends upon the circumstances.

For applications where only a few receivers are to be serviced by a single transmitter, a wire medium is satisfactory. In applications where a great many receivers must be serviced simultaneously by one transmitter or where extremely rugged terrain must be traversed, a wire line is impractical and the space medium is employed.

Depending upon the application of the transmitter and receiver, varying numbers of low and high-frequency amplifier circuits are needed. However, regardless of the nature of the device in which it is contained, each basic circuit operates in its own particular manner.

STAGE-TO-STAGE TRANSFER OF ENERGY

In general, the group of components and circuits associated with one or more electron tubes is known as a stage. With few exceptions, each stage serves only one main function. Thus, the term "amplifier stage", "detector stage", or "oscillator stage" indicates this primary function.

On the other hand, a single function may be performed by a circuit which contains more than one stage. For example, in a radio receiver, the low-frequency amplifier may have two or more stages to increase the level of the audio signal until it is high enough to operate the loudspeaker at the desired volume.

In most devices, the signal energy is passed along from one stage to the next. Modified by each stage in some manner, it finally appears in the desired form at the output of the unit. Thus, it is necessary to provide some type of connection between stages to permit this transfer of energy.

The circuits that connect stages together or connect stages to input and output devices are called **COUPLING NETWORKS** or **COUPLING CIRCUITS**. In addition to coupling energy between stages or from a stage to some device, the coupling circuit must match the impedance of one stage to another or to some input or output device. The coupling network used depends upon the application.

The coupling circuit and load impedances determine the frequency response of the amplifier. Some high-frequency amplifiers employ wideband coupling circuits so a wide range of frequencies can be amplified. However, most high-frequency amplifiers employ tuned coupling and load circuits to provide a restricted frequency response. In this manner, maximum gain is obtained since the amplifier only amplifies a narrow frequency band. This type of amplifier is known as a **TUNED AMPLIFIER**.

H-F COUPLING METHODS

Coupling circuits can be classed as either inductive or capacitive. Capacitive coupling is subclassified into two types according to the types of components employed in shunt with the coupling components. When both the input and output shunt components are resistors, the arrangement is called **RESISTANCE-CAPACITANCE COUPLING**, or simply **RESISTANCE COUPLING**.

If one or both of the shunt components is an inductor, or an LC circuit, the method is known as **IMPEDANCE COUPLING**. In both resistance and impedance coupling, the shunt components are separated electrically and physically so that the only coupling is by means of the series capacitor.

Inductive coupling is subclassified into two types also. The difference is in the method by which the mutual or coupling inductance is obtained. In **TRANSFORMER COUPLING** the energy transfer is a result of the magnetic flux from the primary winding cutting the secondary. In **MUTUAL IMPEDANCE COUPLING** the primary and secondary coils are not coupled magnetically, but a third coil which is part of both circuits provides the coupling.

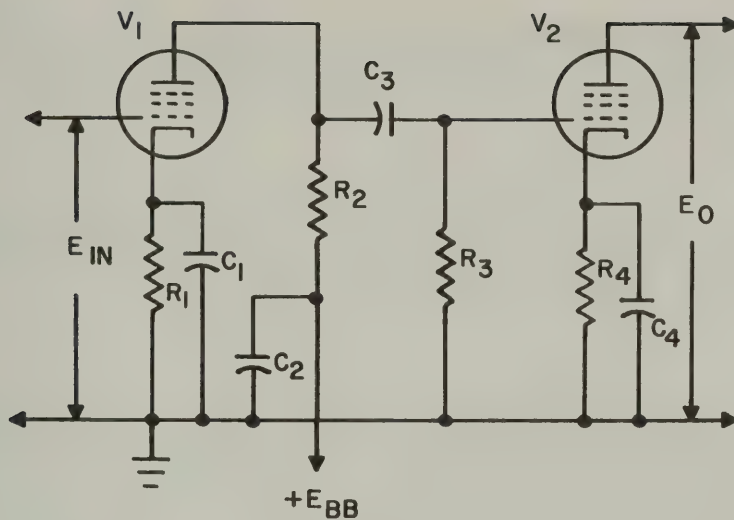


Figure 3

Capacitive Coupling

A resistance-capacitance coupling circuit is illustrated in Figure 3. Here, the signal voltage at the plate of V_1 is transferred to the grid of V_2 by the inter-stage coupling circuit which consists of plate load resistor R_2 , series coupling capacitor C_3 , and grid resistor R_3 . At the signal frequency, bypass capacitor C_2 has a very low reactance so that the $+E_{BB}$ end of resistor R_2 is effectively grounded as pictured in the equivalent circuit of Figure 4.

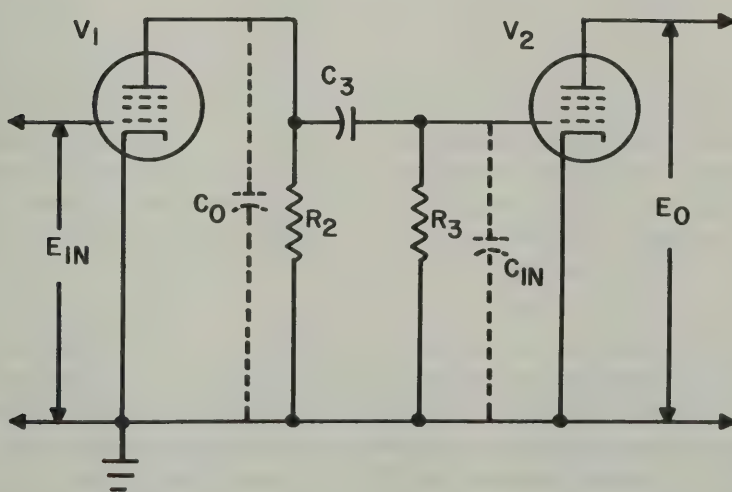


Figure 4

Briefly, the coupling circuit action is as follows. The V_1 plate current variations cause corresponding voltage variations across resistor R_2 . The voltage variations across R_2 cause corresponding variations of the V_1 plate voltage with respect to ground.

Insofar as the signal frequency is concerned, capacitor C_3 and resistor R_3 form a series circuit between the V_1 plate and ground. This circuit forms a highpass filter which is in parallel with R_2 .

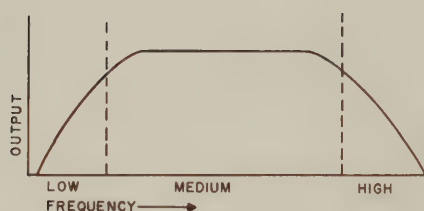


Figure
5

Since the C_3 reactance varies inversely with frequency, the V_1 output actually developed across resistor R_3 depends upon the signal frequency, the C_3 capacitance, and the R_3 resistance. With a given set of components, the C_3 reactance at some low frequency is so high that all of the signal voltage is developed across the capacitor and none appears across the resistor. This condition is pictured by the gradual decrease in output at the low-frequency end in the graph of Figure 5.

As the frequency increases, less voltage is developed across the capacitor and more appears across the resistor. Finally, at some medium frequency, the C_3 reactance is so low that practically all of the signal appears across the resistor.

To couple maximum energy from the V_1 plate to the V_2 grid, the C_3R_3 time constant must be large over the desired frequency range. Typical values of coupling capacitance and grid resistance range from about 10 pF and 10,000 ohms in an FM or television receiver to 500 pF and 1 megohm in a broadcast receiver.

As far as signal transfer is concerned, the series reactance could be reduced to zero by eliminating C_3 and connecting the plate of V_1 directly to the grid of V_2 . Then all the voltage across R_2 is coupled to V_2 . However, for proper operation the V_1 plate must be positive and the V_2 grid negative with respect to their cathodes. C_3 is necessary to allow these electrodes to have the required dc voltages. With low reactance, C_3 offers minimum opposition to the transfer of h-f energy and serves also as a dc blocking capacitor.

It is undesirable that the power supply circuit be allowed to carry the signal currents. Therefore, in Figure 3 capacitor C_2 is chosen to have a low reactance at the operating frequency to prevent h-f currents present at the lower end of R_2 from entering the power supply. Also, capacitors C_1 and C_4 provide long time constants across bias resistors R_1 and R_4 to reduce h-f voltage drops across these resistors.

Typical values of C_2 capacitance range from about $.001 \mu\text{F}$ in FM and television receivers to $.1 \mu\text{F}$ in broadcast receivers. The C_1 and C_4 capacitances range from about $.001 \mu\text{F}$ to $.01 \mu\text{F}$.

From the standpoint of the signal currents, the ideal condition would be that in which the various h-f current paths had zero impedance. Such an arrangement is illustrated in Figure 4, where the three bypass capacitors in Figure 3 are replaced by direct connections to show the signal circuits in simplified form. For simplicity the connections to $B+$ and $B-$ are omitted.

The equivalent circuit diagram of Figure 4 shows that resistors R_2 and R_3 are in shunt across the signal circuit. In any electron tube, INTERELECTRODE CAPACITANCE exists between each pair of electrodes, because electrically the electrodes are conductors separated by a dielectric. Also, in any circuit, stray capacitance exists between the various leads and ground.

In an electron tube circuit, the plate-to-cathode interelectrode capacitance adds to the stray capacitance between the plate lead and ground to form a total known as the OUTPUT CAPACITANCE (C_o). As indicated in Figure 4, the V_1 output capacitance C_o is in parallel with the load resistor R_2 .

Also, in any electron tube circuit, the grid-to-cathode interelectrode capacitance adds to the stray capacitance between the grid lead and ground to form a total called the INPUT CAPACITANCE (C_{IN}) of the stage. Hence, Figure 4 shows the input capacitance C_{IN} in parallel with grid resistor R_3 .

At the lower frequencies, C_o and C_{IN} have extremely high reactances. In this case the parallel impedance of C_o and R_2 is practically equal to R_2 , and the parallel impedance of C_{IN} and R_3 is practically equal to R_3 . Thus, at low frequencies, C_o and C_{IN} have negligible effect on the operation of the circuit. However, the output and input capacitances have decreasing reactance as frequency increases, and at the higher frequencies produce important effects on the signal.

As the frequency increases, the reactances of C_o and C_{IN} decrease. Since the total impedance of the $C_o R_2 R_3 C_{IN}$ combination is less than the equivalent reactance of C_o and C_{IN} in parallel, it may become quite small. The plate resistance of tube V_1 does not change appreciably with frequency. Hence, more of the signal voltage appears across V_1 and less across the load. The result of this action is a decrease of output as shown by the curve of Figure 5 until, at some high frequency, the C_o and C_{IN} reactances are so low that they short-circuit the signal and reduce the output to zero.

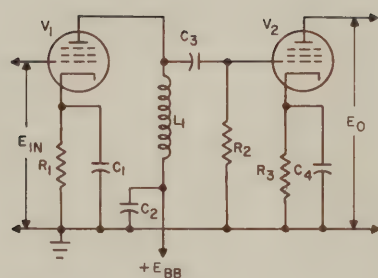


Figure 6

One method of improving the high-frequency response of an amplifier is to lower the R_2 resistance. Of course, this change lowers the output at the medium and low frequencies because it lowers the plate load impedance. But, what is more important, it raises the frequency at which the shunt capacitance reactance becomes an appreciable part of the load impedance.

An impedance-coupled circuit is shown in Figure 6. Here, inductor L_1 forms the shunt component at the input of the coupling circuit, C_3 is the series coupling capacitor, and R_2 is the shunt component across the coupling circuit output. Across the signal circuit, high reactance and resistance are presented by L_1 and R_2 , respectively, while the series arm C_3 offers very low reactance at high signal frequencies. Thus, as in Figure 3, the coupling circuit in Figure 6 is a pi-type high-pass filter. Again, C_3 serves as both coupling and dc blocking capacitor between the V_1 plate and the V_2 grid.

In Figure 6, bypass capacitors C_1 , C_2 , and C_4 connect the cathodes and lower end of L_1 to ground so far as signal currents are concerned. Thus, for this circuit, the signal circuits are like those of Figure 4. In Figure 6 the h-f signal voltage appears across L_1 , and except for the small loss across C_3 this voltage is also impressed across R_2 in the grid circuit of V_2 .

Interelectrode and stray capacitances are present in any electron tube circuit, and in Figure 6 the distributed capacitance of L_1 increases the output capacitance of the V_1 stage. This capacitance and L_1 form a parallel resonant circuit at some high frequency. Since the impedance of a parallel resonant circuit rises as the frequency approaches resonance, this type of coupling circuit has high shunt impedance at the higher signal frequencies. Thus, the inductive load develops a relatively high output voltage which is coupled to the grid of V_2 . If the value of L_1 is chosen to resonate with the shunt capacity, high gain at some high frequency can be obtained. In most circuits, however, an actual tuning capacitor is included to resonate with L_1 at this particular frequency.

A second advantage of the plate load inductor is that, while having high reactance at the signal frequencies, it has a very low dc reactance. Therefore, the IR drop is low compared with that of a resistance load, which makes it unnecessary for the $+E_{BB}$ voltage to be much higher than the operating voltage required for the tube plate. The lower IR drop also reduces the power dissipated by the load resistor.

The arrangement of Figure 3 may be employed when a comparatively uniform response over a relatively wide frequency range is desired. However, for most high-frequency receiver and transmitter applications a high amplification over a comparatively narrow band is required. This can easily be accomplished

by adding a capacitor in parallel with L_1 in Figure 6. Figure 7 shows such a circuit employing a resonant plate load.

Connected directly in series with the plate, capacitor C_3 and inductor L_1 form the parallel resonant circuit that is the load impedance for tube V_1 . C_4 is the coupling capacitor through which the h-f energy in the V_1 plate circuit is transferred to the grid of the following tube. C_2 is a bypass capacitor to decouple the h-f energy from the power supply.

Grid bias voltage for V_1 is provided by cathode resistor R_1 and capacitor C_1 . Shown in broken lines, capacitor C_0 represents the V_1 output capacitance, and capacitor C_{IN} represents the input capacitance of the following tube. At high frequencies the reactances of C_2 and C_4 are extremely low so that, in effect, C_0 and C_{IN} are connected in parallel with L_1 . Thus, C_0 and C_{IN} add to C_3 to increase the total capacitance which tunes L_1 .

Applied to the V_1 grid, input signal E_{IN} is amplified and appears across the grid resistor R_2 as output voltage E_O . The voltage amplification of a stage may be expressed as:

$$A_V = \frac{E_O}{E_{IN}} \quad (2)$$

where

A_V is the voltage amplification,

E_O is the output voltage,

E_{IN} is the input voltage.

When the amplifier tube is a pentode, the amplification may be stated as:

$$A_V = g_m \times Z_L \quad (3)$$

where

A_V is the voltage amplification,

g_m is the tube transconductance,

Z_L is the load impedance.

As an example, suppose the plate circuit impedance is 100 k Ω and the g_m of the amplifier tube is 7,000 μ mhos. The stage gain is therefore equal to

$$A_V = g_m \times Z_L = 7,000 \times 10^{-6} \times 1 \times 10^5 = 7,000 \times 10^{-1} = 700.$$

(Remember; the transconductance, g_m , must be expressed in basic units, mhos)

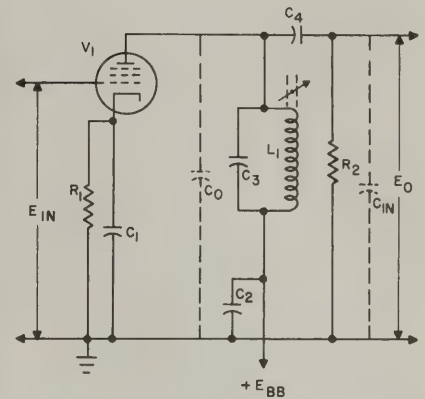


Figure
7

VACUUM TUBE HIGH FREQUENCY AMPLIFIERS

Since the coupling circuit contains one tuned circuit, the amplifier is called a **SINGLE-TUNED AMPLIFIER**. In the tuned circuit, the inductive reactance is lower for frequencies below resonance and the capacitive reactance is lower for frequencies above resonance. Because these reactances vary with frequency, the impedance Z_L of the tank is maximum at the resonant frequency of the tuned circuit, and falls off for frequencies above and below resonance. The tube transconductance remains constant; therefore, amplification (A_v) varies with frequency and is maximum at resonance.

The curves of Figure 8 show how the amplification varies over a range of frequencies in the case of two different amplifiers. In both examples the tuned circuit is resonant at 1000 kHz. Thus, the amplification is maximum at resonance and lower at frequencies above and below.

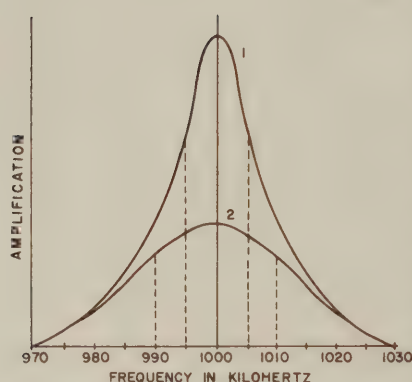


Figure
8

Figure 8 shows how the response of the amplifier varies with the Q of the tuned circuit. In Figure 8, curve 1 represents the response of an amplifier in which the tuned circuit has a relatively high Q , while curve 2 represents the response of a low- Q circuit. Due to the shunting effects of grid resistor R_2 and the internal plate resistance of tube V_1 in Figure 7, the effective Q of the tuned circuit always is less than the Q of the coil. Therefore, the circuit amplification is lower than indicated by a curve showing coil Q only.

Most signals include a range or band of frequencies. But the curves of Figure 8 indicate maximum amplification for only the resonant frequency. In



In a communications receiver, like the one shown above, double-tuned transformer coupled amplifiers are employed in the intermediate amplifier stages.
Courtesy R. L. Drake Co.

practice, it is generally considered that all frequencies with magnitudes equal to or greater than 70.7% of maximum are amplified properly. These frequencies are said to be "passed" by the amplifier. They cover a range that is known as the passband. Depending upon the width of the passband, in a given application an amplifier may be considered to have narrow, medium, or wide frequency response or **BANDPASS**.

The bandpass of a tuned amplifier depends upon the Q of the tuned circuit. The broken lines in Figure 8 indicate a bandpass from 995 to 1005 kHz for a high- Q circuit with the response of curve 1, and a bandpass from 990 to 1010 kHz for a lower- Q circuit with the response of curve 2. Thus, the lower Q permits a broader bandpass but results in a lower amplification.

The amplification or gain in a high-frequency amplifier is also influenced by the inductive and capacitive reactances in the tuned circuit. The gain is high when both of these reactances are high. The coil must have large inductance to provide high reactance at a given frequency, but for high capacitive reactance the circuit capacitance must be small. Therefore, high gain is obtained when the **TUNED CIRCUIT HAS A LARGE L TO C RATIO**.

When very wide bandpass is desired, the Q of the circuit must be low, and the gain of the stage is reduced. To compensate for the loss, the L to C ratio must be as high as possible to maintain the gain at a practical level. When carried to extremes, the result is an amplifier in which no actual capacitor is connected across the inductor, and the circuit is resonated by the interelectrode and stray capacitances represented by C_o and C_{IN} . Hence, it becomes necessary to vary the inductance to tune the tank circuit. In Figure 7, this is accomplished with a movable powdered iron core in the inductor as indicated by the dotted lines and arrow above the inductor symbol. Moving this core varies the permeability of the magnetic circuit, and this variation results in a corresponding change in inductance. Therefore, the movable core tunes the tank circuit to resonance at the signal frequency.

Inductive Coupling

In the transformer-coupled circuit of Figure 9, inductors L_1 and L_2 are located close together physically, and form the primary and secondary, respectively, of an h-f transformer.

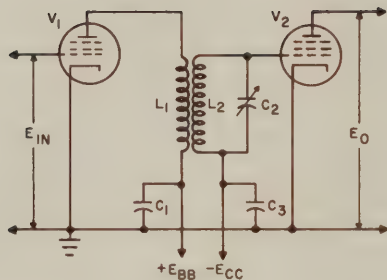


Figure 9

Iron cores similar to those employed in low-frequency transformers would cause excessive losses at high frequencies, and so are unsuitable. Therefore, most h-f transformers employ either powdered iron or air cores. However, due to the relatively high reluctance of these cores, the primary current is unable to produce a strong magnetic field. Therefore, the induction between the primary and secondary is relatively low. Hence, the degree of coupling becomes an important factor in determining the response which the transformer has to a range of signal frequencies.

In radio receivers the coupling is adjusted to the desired value by positioning the coils properly and then sealing them so that they remain fixed. In some applications, which require a coupling that can be varied at will, one coil is made movable so that its position with respect to the other may be changed easily.

Due to transformer action, the V_1 plate circuit appears to contain the tuned circuit impedance, and since the amplifier is single tuned, the frequency response is as shown by the curves of Figure 8. As before, curve 1 represents the response when the tuned circuit has a relatively high Q , and curve 2 represents the response of a circuit with lower Q .

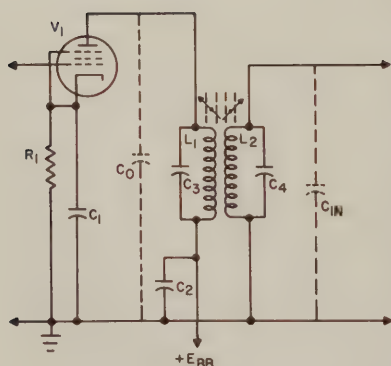


Figure 10

The response curves of the circuits of Figures 7 and 9 are sharply pointed when relatively high- Q resonant circuits are employed. Although this type of response is satisfactory for many applications, there are others such as in broadcast and FM radio receivers which require a steep-sided, relatively flat-topped response characteristic. The flat-topped response provides substantially constant amplification for all the desired signal frequencies, and at the same time discriminates sharply against frequencies outside the desired band.

A steep-sided, flat-topped response may be obtained by tuning both the primary and secondary of the coupling transformer, as shown in Figure 10. The amplifier then becomes a **DOUBLE-TUNED AMPLIFIER**. Capacitor C_3 , in parallel with output capacitance C_0 of V_1 tunes transformer primary L_1 , while secondary L_2 is tuned by the parallel combination of C_4 and C_{IN} .

With two tuned circuits resonant to the same frequency, the frequency response depends on the degree of coupling between the coils. When the degree

The following Practice Exercise questions cover the subjects which you have just studied. They are:

WAVELENGTHS

FREQUENCY SPECTRUM

HIGH FREQUENCY ELECTRONIC EQUIPMENT

STAGE-TO-STAGE TRANSFER OF ENERGY

H-F COUPLING METHODS

CAPACITIVE COUPLING

1. Wavelength is (directly) (inversely) proportional to frequency.
2. What is the wavelength in meters of a 15 MHz radio wave?
3. What is the name given to the group of components and circuits associated with one or more electron tubes?
4. Describe what is meant by a tuned amplifier.
5. How does resistance coupling differ from impedance coupling?
6. How does transformer coupling differ from mutual impedance coupling?
7. The RC coupling circuit in the circuit of Figure 3 is a (high pass) (low pass) filter.
8. For proper circuit operation, C_3 and R_3 in Figure 3 should have a short time constant. True or False?
9. What is the function of C_2 in Figure 3?
10. How do the input and output capacities shown in Figure 4 affect circuit operation?
11. What is the effect of lowering the value of R_2 in Figure 3?
12. What happens to the gain in the circuit of Figure 6 at the frequency where L_1 resonates with the stray capacity?
13. Why is the low dc resistance of L_1 in Figure 6 an advantage over a simple load resistor?
14. Is the expression for the voltage gain of a pentode amplifier different from that for a triode amplifier?

15. Why is the circuit of Figure 7 called a single-tuned amplifier?
16. What affect does the Q of the tuned circuit have on the response of the amplifier of Figure 7?
17. How does the impedance of the L_1C_3 plate load in Figure 7 vary with the LC ratio? How does this affect the gain?
18. The bandpass of a tuned amplifier is considered to be all frequencies with amplitudes equal to or greater than _____% of maximum.

of coupling is small, the interaction between the primary and secondary is low and the amplification curve resembles that of ordinary resonance, as shown by curve 1 in Figure 11. Also, due to the small coupling the output voltage is low so that the amplification is small.

As the coupling between the coils is increased, the amplification increases to a maximum. This point is called "CRITICAL COUPLING". If the coupling is increased beyond the critical point, the response curve takes on a double-humped appearance as shown by curve 2 in Figure 11. A further increase of coupling does not increase the amplification, but causes the two humps to spread apart and the valley between them to deepen. Again, the shape of the amplification curve depends on the Q's of the tuned circuits, being higher and steeper sided for higher-Q circuits.

Thus, an amplifier employing the circuit of Figure 10 may be adjusted to provide essentially uniform response at all frequencies within a desired band, but little response to frequencies beyond the band limits. This poor response beyond the desired band is described by saying that the circuit has good SKIRT SELECTIVITY.

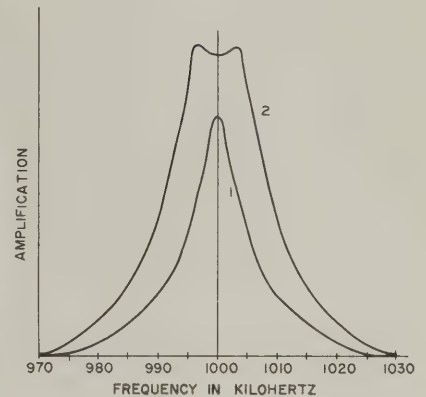


Figure 11

In the circuits of Figures 3, 6, and 7, grid bias voltages are obtained by means of the voltage drop due to plate current in the cathode resistor. In the case of V_2 in Figure 9, the cathode connects to ground, and the negative direct voltage for the grid is obtained from the fixed bias supply $-E_{CC}$. In this case, L_2 cannot be connected directly to ground, but an h-f current path to ground is provided by decoupling capacitor C_3 . Therefore, the bias supply is in series with the grid-cathode dc circuit.

Although illustrated only in Figure 9, the fixed bias arrangement may be employed in other types of coupling circuits, such as those of Figures 3 and 6. Also, in Figure 9, cathode resistors can be employed to provide the grid bias for V_1 and V_2 if desired. Then L_2 would be connected to ground, and C_3 placed in parallel with the cathode resistor like C_1 in Figure 6. The V_1 grid and cathode circuits would be completed in a similar manner. In some circuit designs, combinations of bias methods are used.

In the circuit of Figure 10, coupling is through the magnetic field developed by the primary. Coupling between two inductors can also be accomplished by employing a third inductor to "mutually" couple the inductors. Figure 12 shows two electron tubes connected by a mutual impedance coupling circuit. Here coils L_1 and L_2 are physically situated so there is no magnetic coupling between them, and the mutual inductance is provided by coil L_3 . The circuit arrangement is equivalent to a transformer in which L_1 and L_3 form the primary inductance, while L_2 and L_3 constitute the secondary inductance. In this circuit, C_2 functions only as a dc blocking capacitor.

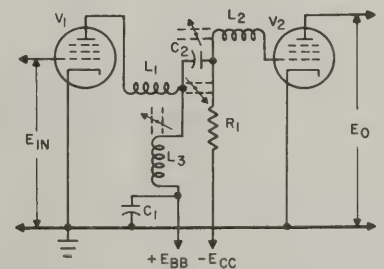


Figure 12

From the plate of V_1 to ground, the signal circuit includes L_1 , L_3 , and C_1 . The V_2 grid-to-ground signal circuit contains L_2 , C_2 , L_3 , and C_1 . Briefly, the coupling action is as follows.

In the V_1 plate circuit, the h-f current in L_3 produces a voltage drop which causes a corresponding h-f current in the V_2 grid circuit. This grid circuit current produces voltage drops across both L_2 and L_3 , and the sum of these drops is the signal voltage applied between the grid and cathode of V_2 . Compared to those across the coils, the voltage drops across C_1 and C_2 are negligible at the signal frequency.

In the primary circuit, L_1 and L_3 are resonated at the signal frequency by the output capacitance of V_1 , while the input capacitance of V_2 resonates L_2 and L_3 at the same frequency. The primary and secondary circuits are tuned to resonance by adjustment of the movable iron cores of L_1 and L_2 , respectively. Adjusting the position of the core in L_3 changes the degree of coupling between stages.

In most cases, inductive coupling requires that the stages be close together. However, the physical size of the circuit components, mode of operation, or other considerations may make it necessary to place the stages some distance apart. To permit separated stages to employ inductive coupling, the circuit arrangement of Figure 13 is used.

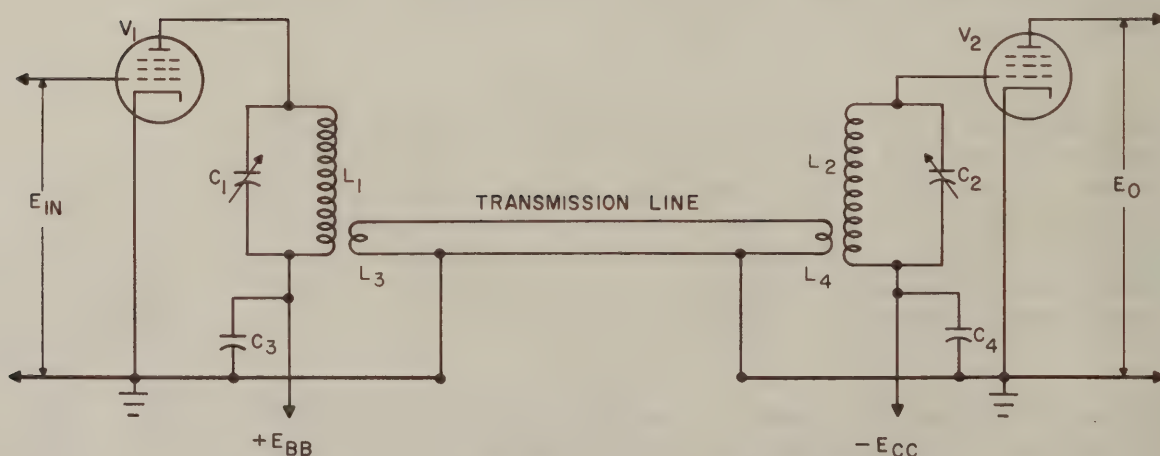


Figure 13

Known as **LINK COUPLING**, this system employs a transmission line to carry the h-f energy from one circuit to the other. The "link" is a line with

small windings of one or two turns at each end. As shown, link winding L_3 is inductively coupled to coil L_1 in the plate tank circuit of V_1 .

A varying flux produced by h-f current in L_1 induces an emf in L_3 . The induced emf produces corresponding h-f currents in L_3 , the transmission line, and L_4 . Caused by this current, a varying flux about L_4 induces an emf in coil L_2 . With tuned circuit L_2C_2 resonated at the signal frequency, the emf causes a large h-f current in the circuit, which results in a large voltage across L_2 and C_2 . This signal voltage is applied to the grid of V_2 .

HIGH-FREQUENCY AMPLIFIER TUBES

When a triode tube is employed in a high-frequency amplifier, the relatively large grid-to-plate capacitance provides a low-impedance path through which h-f energy from the plate is fed back to the grid circuit. If the phase and magnitude of the feedback voltage are correct, it reinforces the original grid signal and causes unstable operation or oscillations.

The tetrode tube was developed primarily to minimize the grid-to-plate capacitance, and thus prevent oscillation. Oscillation is a condition where an amplifier generates an output signal without any input signal. The screen grid functions as a shield between the grid and plate, thereby reducing the capacitance between these two electrodes to a very low value. In pentode type tubes, an additional electrode called the suppressor grid is placed between the screen grid and plate to reduce the secondary emission effects.

Because of their extremely low grid-to-plate capacitance, tetrode and pentode tubes are widely used as voltage amplifiers in many types of high-frequency electronic equipment. However, the power-handling capabilities of tetrodes and pentodes are limited, and triodes are used frequently in high-power applications.

H-F VOLTAGE AMPLIFIERS

Generally speaking, a voltage amplifier is one which provides a high voltage gain. Therefore, a small signal input voltage produces a large signal output voltage. This type of amplifier takes practically no power from the preceding stage. It also delivers practically no signal power to the following stage. Thus, from a power standpoint it is very inefficient. A voltage amplifier is thus used where power is not required, such as in the i-f and r-f stages of radio, television, and communications receivers. Tetrodes and pentodes are usually used as voltage amplifiers due to their high amplification factor.

VACUUM TUBE HIGH FREQUENCY AMPLIFIERS



A communications transceiver, like the one shown above, employs voltage amplifiers in the receiver section and power amplifiers in the transmitter section.

Courtesy Heath, Inc.

An h-f voltage amplifier may be either tuned or untuned. In the tuned amplifier a parallel resonant circuit provides the high plate load impedance. In an untuned amplifier, the plate load is not resonant at the signal frequency. Tuned h-f voltage amplifiers almost always employ pentode tubes, because of the high voltage amplification obtained, and because the very low grid-to-plate capacitance reduces feedback.

H-F POWER AMPLIFIERS

In many electronic systems high-frequency signals at large power levels are required. Transmitters and dielectric heaters are typical examples. To obtain these high power levels, triodes are usually used. In cases where the power requirements are not high, tetrodes and pentodes may be used but their efficiency drops off rapidly above 100 MHz.

The power delivered to an amplifier by the power supply is the product of the instantaneous current and voltage. Part of this power is delivered to the load and represents useful output. The remainder is dissipated in the form of heat at the plate of the tube. At any instant, the total power divides between the load and tube in proportion to the voltage drops across these parts of the circuit. Thus, the power dissipated by the tube is equal to the product of the instantaneous plate current and voltage. The average input, output, and plate loss powers are found by averaging the instantaneous values over a complete cycle.

To obtain maximum power output from a given amplifier with given power input, the plate loss must be reduced to a minimum. This can be done by permitting plate current only when the plate voltage is low, as shown in Figure 14, where curve A indicates a sinewave plate voltage e_b , which is the result of the resonant tank circuit in the plate load, and curve B shows that plate current i_b is in the form of pulses which occur only when the plate voltage is near its minimum.

Both Class B and Class C amplifiers operate with these pulse type currents. In Class B amplifiers, the grid bias is approximately equal to the plate current cutoff value so that the plate current is in the form of pulses lasting for approximately half of each cycle of the input signal. In Class C amplifiers the grid is biased beyond cutoff so that the plate current pulse lasts for appreciably less than half of the input cycle.

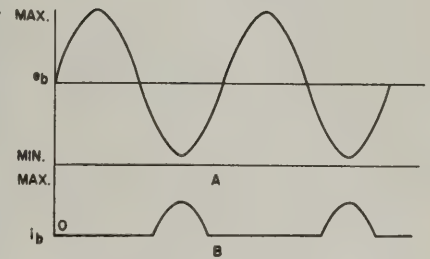


Figure 14

At the plate of the tube, the power loss is equal to the average of $e_b i_b$. Since i_b is zero at all times except when e_b is near its minimum, the power loss at the plate is low, and a large amount of power is delivered to the load. The PLATE CIRCUIT EFFICIENCY of an amplifier is the ratio of the ac power in the load to the dc power delivered from the power supply. Usually, this ratio is stated as the percent of efficiency thus:

$$\% \text{ Eff.} = \frac{P_o}{P_{IN}} \times 100, \quad (4)$$

where P_o is the ac power in the load, and P_{IN} is the dc power delivered by the supply that is, $E_b \times I_b$. Thus, when the plate loss is low, the power in the load is high, and the efficiency is high.

As an example, suppose that an amplifier delivers 100 watts to a load and the plate voltage is 800 volts and the average plate current is 150 mA. The dc input from the power supply is equal to:

$$P_{IN} = E_b \times I_B = 800 \text{ volts} \times .15\text{A} = 120 \text{ watts.}$$

The efficiency is therefore equal to:

$$\% \text{ Eff.} = \frac{P_o}{P_{IN}} \times 100 = \frac{100 \text{ watts}}{120 \text{ watts}} \times 100 = 83\% \text{ (approx.)}$$

Because their shorter plate current pulses make them the most efficient type, Class C amplifiers are employed when waveform distortion is permissible.

When this type of amplifier is in operation, the high negative grid bias allows plate current in short pulses that do not have a sine waveform, even though the grid is excited by a sinewave voltage. However, these current pulses store up a charge in the tuning capacitor and as the plate current drops to zero, the capacitor discharges through the inductor and stores the energy in the magnetic field.

A constant interchange of energy thus occurs between the inductor and capacitor. As a result, a sinusoidal "oscillating" current is set up in the tuned circuit. To maintain this oscillating current, the tuning capacitor is adjusted so that the resonant frequency of the tuned circuit is the same as that at which the plate current pulses occur.

Frequently, this oscillating current action is referred to as the FLYWHEEL EFFECT, an example of which is found in the balance wheel of a watch. Pivoted between jeweled bearings to reduce friction, the balance wheel oscillates back and forth at a rate that depends on the mass of the wheel and the strength of the hair-spring. Each time the watch ticks, the mainspring supplies a little push to the wheel by way of the escapement mechanism.

Due to the momentary pushes repeated at just the right instants, energy which is lost due to friction is replaced, the balance wheel continues to oscillate at the same rate, and the watch keeps running. Without the little pushes from the mainspring, the friction would gradually absorb all of the energy and the balance wheel would slow to a stop.

The same conditions exist in the plate tank circuit. A single pulse of energy applied to the circuit causes a circulating current at the resonant frequency. Unless additional energy is supplied from time to time, the current gradually falls to zero.

On the other hand, the circulating current is sustained if additional energy is supplied periodically and at the proper instants. If the energy pulses are equal, the amplitude of the circulating current is constant. Thus, although the plate current is in short pulses, the voltage developed across the tuned circuit has a sine waveform.

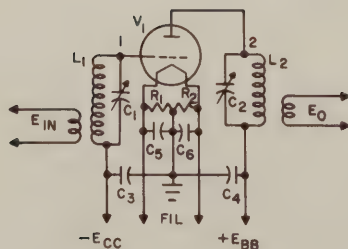


Figure
15

Exactly the same circuits may be employed for both Class B and Class C amplifiers, the only difference being the value of the grid bias. In the typical triode h-f amplifier stage of Figure 15, the input signal is inductively coupled into tuned circuit L_1C_1 , from which it is impressed between the grid and cathode of the tube. The amplified signal is developed across tuned circuit L_2C_2 from which it is inductively coupled to the output circuit. Capacitors

C_3 and C_4 provide low reactance paths at the signal frequencies, thus decoupling the h-f energy from the plate and bias power supplies.

With a directly-heated emitter or tube filament and the customary ac supply, the filament voltage must be decoupled from the grid and plate circuits to eliminate hum. As the filament is in series in both the grid and plate circuits, the effects of the changes of voltage across it can be greatly reduced by connecting to its electric center.

As illustrated in Figure 15, one common method of providing this center connection is to connect equal value resistors, R_1 and R_2 , in series across the filament. Under these conditions the junction between the resistors is at the same potential as the electric center of the filament and is connected directly to ground. Connected in parallel with the resistors, capacitors C_5 and C_6 offer low impedance at the signal frequencies; in effect, they short the filament voltage insofar as the grid and plate signal circuits are concerned.

Grid Driving Power

To obtain maximum power output from a Class B or Class C amplifier stage, the signal input voltage must overcome the negative bias and drive the grid positive during a portion of each cycle. As indicated by curve B of Figure 14, the plate current occurs in pulses and reaches maximum when e_b is minimum. Remember, maximum plate current causes maximum voltage drop across the load, and thus reduces the plate voltage to a minimum. Therefore, the power output is determined largely by the amount the grid is driven positive.

However, there are certain limitations because, when driven positive, the grid carries current and the resulting power must be dissipated in the form of heat. To prevent excessive heat from damaging the grid structure, the grid power must be kept below the manufacturer's recommended value.

Also, when grid current exists, power is absorbed by the resistance of the grid circuit, including the bias supply. For proper operation of Class B or Class C amplifiers, the driving source must have sufficient power capacity to overcome the circuit losses.

Push-Pull Amplifiers

Although the explanations in this lesson are based mainly on single tube amplifier circuits, high-frequency amplifiers may also be operated in push-pull. As in low-frequency amplifiers, this type of operation results in twice the power output of one tube and a reduction of even harmonic voltages. This latter feature is very important because radio and television transmitting stations must keep within a federal law limiting harmonic radiation.

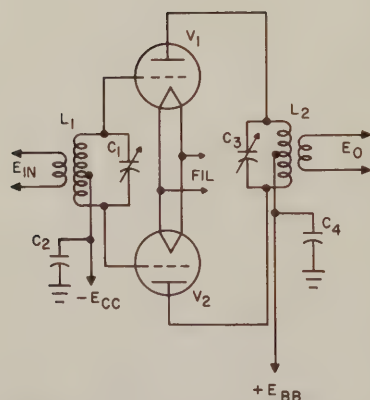


Figure
16

Figure 16 shows the circuit diagram of an h-f push-pull amplifier. The required 180° out of phase grid signals are obtained by using a center tapped tank circuit in the grids of V_1 and V_2 . C_1 tunes the secondary of L_1 to the desired operating frequency. Fixed bias is employed and C_2 bypasses the bias supply. Either Class B or Class C operation can be obtained by selecting the proper value for $-E_{CC}$. C_3 tunes the center tapped primary of L_2 , the plate tank circuit, to the desired operating frequency. Although not shown, a center tapped grounding arrangement, like that employed in Figure 15, would be used for the filament circuit. C_4 bypasses the plate voltage supply.

FREQUENCY MULTIPLIERS

In high frequency equipment it is often necessary to increase a signal frequency. That is, we may wish to double or triple a signal frequency. As an

The following Practice Exercise questions cover the subjects which you have just studied. They are:

H-F COUPLING METHODS

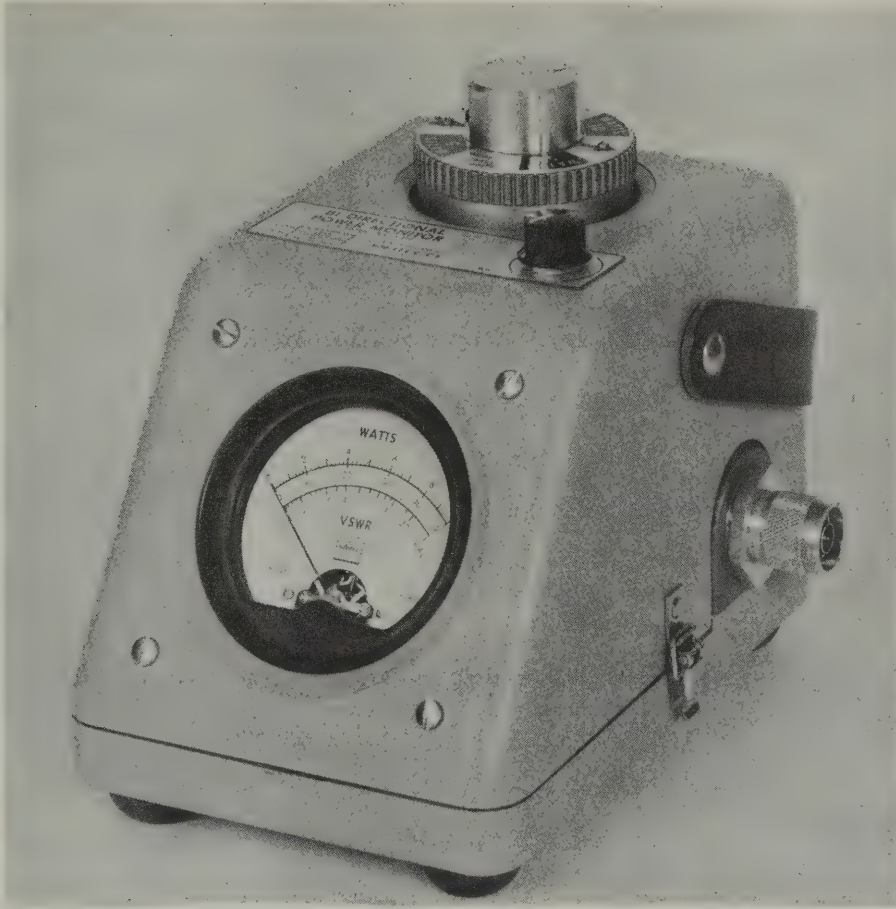
INDUCTIVE COUPLING

HIGH-FREQUENCY AMPLIFIER TUBES

H-F VOLTAGE AMPLIFIERS

H-F POWER AMPLIFIERS

19. The coupling transformer in an h-f amplifier generally employs a laminated iron core. True or False?
20. What happens to the bandwidth of a double-tuned amplifier if the coupling is increased beyond the critical point?
21. When the sides of a response curve are steep sided, is the skirt selectivity high or low?
22. How does mutual impedance coupling differ from transformer coupling?
23. In the circuit of Figure 12, which inductor is the "mutual impedance"?
24. What is the advantage of link coupling?
25. What is the advantage of using a tetrode as compared to a triode in an h-f amplifier?
26. The i-f and r-f stages in a television receiver would be classed as voltage amplifiers. True or False?
27. What is meant by the efficiency of an amplifier?
28. Suppose a h-f power amplifier is supplying 200 watts to a load. If the plate voltage is 750 volts and the plate current is 500 mA, what is the efficiency?
29. How is a sinewave current produced in the plate tank circuit of a Class C amplifier?
30. A Class C amplifier yields a high efficiency since plate current pulses occur when plate voltage is (high) (low).
31. What changes in a Class B amplifier are necessary to operate it Class C?
32. When using a directly heated tube in an h-f amplifier it is common practice to ground either side of the filament. True or False?



The unit shown above can be used to measure the r-f power delivered to a load from an h-f power amplifier.

*Courtesy Sierra Electronic Operation
Philco-Ford Corp.*

example we may have a 5 MHz signal and desire a 10 MHz signal. The 10 MHz signal can be derived from the 5 MHz signal by employing a frequency multiplier.

A **FREQUENCY MULTIPLIER** is a Class C amplifier operated with its plate circuit tuned to a harmonic of the grid circuit frequency. Usually, the multiplying factor is 2, 3, or 4, although higher factors are employed occasionally. Because the Class C amplifier is biased well beyond cutoff, plate current is in the form of short duration pulses which are rich in harmonics. When the plate circuit is tuned properly for multiplier operation, it presents a high impedance to the desired harmonic, but not to the fundamental or the other harmonic frequencies. Consequently, a large output is produced at the desired harmonic frequency.

Figure 17 shows the schematic diagram of a frequency multiplier stage. This circuit is nothing more than a grid leak biased h-f amplifier operating Class C.

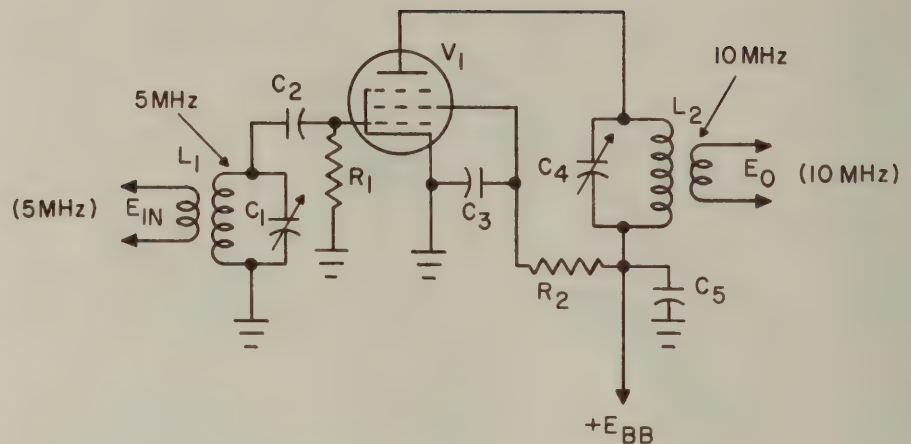


Figure 17

As mentioned earlier, Class C operation produces a distorted signal since the plate current is in the form of pulses. This produces a signal in the plate tank circuit that is rich in harmonics. The desired harmonic is selected by tuning the plate tank circuit to the desired frequency. In the circuit of Figure 17, the input signal frequency is 5 MHz. C_1 tunes the secondary of L_1 to 5 MHz. The plate tank circuit, C_4 and L_2 , is tuned to the second harmonic (10 MHz) of the input signal. The plate tank circuit thus offers a high impedance to this frequency but a low impedance to the fundamental and other harmonics. As a result, only the second harmonic, 10 MHz, is coupled to the output. When the output is tuned to twice the input, as in Figure 17, the circuit is called a DOUBLER. If the output circuit is tuned to three times the input, the circuit is called a TRIPLER and so on.

To obtain good efficiency in frequency multipliers, the plate current pulse should last for less than half a complete cycle at the harmonic frequency. Therefore, the grid bias voltage must be increased until the necessary plate current pulse duration is obtained. At the same time, the h-f input voltage must be increased sufficiently to drive the tube to the same peak plate current. Due to the short duration of plate current, the average plate current is reduced and the power output from the frequency multiplier is not as great as when the plate and grid circuits are tuned to the same frequency.

NEUTRALIZATION

When the grid and plate circuit of a high frequency amplifier are tuned to the same frequency, oscillation can occur. Oscillation is a condition where the amplifier circuit develops an output signal without any input. The frequency of this signal is generally the same as the resonant frequency of the

tuned circuits. Oscillation results due to the fact that the grid-to-plate capacity couples energy from the plate circuit back into the grid circuit.

Although the condition of oscillation is desirable in an oscillator circuit, it is not wanted in an amplifier. An amplifier must only amplify the input signal. It must not develop any signals of its own. The factor which causes oscillation in an amplifier, the grid-to-plate capacity, cannot be eliminated. However, it is possible to neutralize the effects of the grid to plate capacity by a circuit referred to as a "neutralization circuit". The neutralizing circuit is basically a feedback circuit which applies a signal to the grid circuit of the same magnitude but opposite phase as that fed back through the grid-to-plate capacity. This "bucking" voltage may be obtained by connecting an adjustable capacitor between the grid and plate circuits to provide the proper value and polarity signal. This process of nullifying the voltage fed back through the interelectrode capacitance of an amplifier tube by providing an equal voltage of opposite polarity is called **NEUTRALIZATION**.

Two basic methods of neutralization in common use are: (1) Plate or Hazeltine, and (2) Grid or Rice. A method applied to push-pull amplifiers is a form of plate neutralization and is known as cross neutralization or push-pull neutralization.

Plate Neutralization

Figure 18 shows the schematic diagram of a high-frequency amplifier which employs **PLATE NEUTRALIZATION**. This arrangement is also referred to as the Hazeltine method. Except for the neutralizing circuit, the circuit of Figure 18 is the same as that of Figure 15, and the corresponding parts perform the same functions. The plate coil is center tapped to form two equal sections indicated as L_{2A} and L_{2B} . The center tap is connected to E_{BB} , and only the L_{2A} section carries the dc plate current. However, variable capacitor C_2 is connected across the entire coil, and when the circuit is tuned to resonance at the signal frequency, the circulating tank current is carried by the entire coil. With the center tap of L_2 as the reference point, the h-f voltages at the opposite ends of the complete coil are always 180° out of phase with each other. From point P at the upper end of L_{2A} , there is undesirable feedback of voltage through interelectrode capacitance C_{gp} to the grid. To neutralize the circuit, voltage is fed back to the grid through neutralizing capacitor C_N from point N at the lower end of L_{2B} .

Capacitor C_N is adjusted so that the energy it feeds back to the grid exactly equals that fed back through C_{gp} . Under these conditions, both feedback voltages have equal amplitude and, because of the 180° phase relationship, they cancel. Thus, the only energy left across the grid-cathode circuit is that which is induced in coil L_1 by signal current in the r-f input coil.

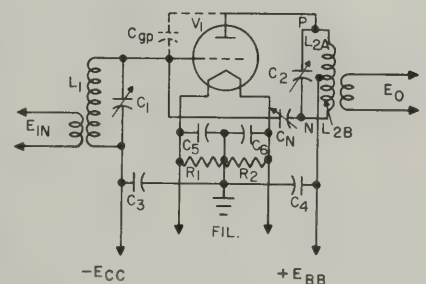


Figure 18

Grid Neutralization

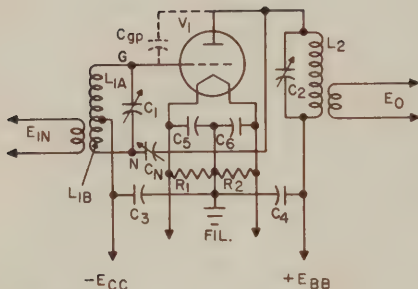


Figure 19

The schematic diagram of a high-frequency triode amplifier employing **GRID NEUTRALIZATION** is shown in Figure 19. This arrangement is also referred to as the **RICE** neutralization method. Again, except for the neutralizing arrangements, the circuit of Figure 19 is the same as that of Figure 15, and the corresponding parts perform the same functions. In this system, grid coil L_1 is center tapped to form two equal sections indicated as L_{1A} and L_{1B} . The upper or G end of the coil is connected to the grid in the usual manner. At the lower end, point N is connected through neutralizing capacitor C_N to the V_1 plate.

With this circuit arrangement, h-f energy from the plate is fed back to the grid circuit through both C_{gp} and C_N . Although both feedback voltages have the same polarity, they are applied to the opposite ends of the grid coil and produce currents that have opposing effects in the two parts of the coil. When neutralizing capacitor C_N is adjusted properly, the effects of the two currents cancel and oscillation is prevented.

THE GROUNDED-GRID AMPLIFIER

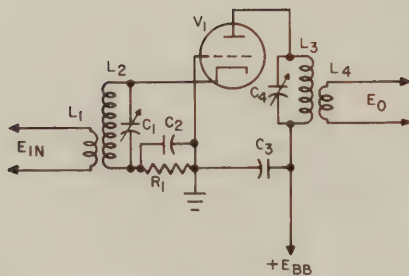


Figure 20

As explained previously, tetrodes and pentodes serve very well in low-power, high-frequency amplifiers without the need for external neutralization. However, at the ultra high frequencies, the shielding effect of the screen grid is reduced and neutralization is necessary to obtain stable operation.

Used primarily to overcome the need for neutralization, the **GROUNDED-GRID AMPLIFIER** employs the circuit of Figure 20. Input circuit L_2C_1 is in the cathode circuit in series with the cathode bias circuit consisting of resistor R_1 and capacitor C_2 . The grid is connected directly to ground and the output is taken from the plate circuit as explained for Figure 15.

In the ordinary grounded-cathode amplifier, the h-f signal is impressed between the grid and cathode, and the cathode is grounded either directly or through a biasing arrangement. Thus, the signal voltages developed across the input circuit cause the grid potential to vary with respect to the grounded cathode.

In the circuit of Figure 20, the input grid-cathode circuit is practically the same except that the grid instead of the cathode is grounded. The signal voltages developed across tuned circuit L_2C_1 cause the cathode potential to vary with respect to the grounded grid. Therefore, with equal input signals, the grid-to-cathode voltages of both the grounded-cathode and grounded-grid

amplifier circuits are equal, and the grids exercise the same control on the plate current.

A severe disadvantage of the grounded-grid amplifier is that the ac plate current must pass through the input circuit. Unless this input circuit resistance is very low, the resulting degeneration will reduce the gain to below that obtainable from the same tube used in a grounded-cathode circuit.

Energy fed back through the interelectrode capacitance from the plate to the grounded-grid does not affect the input circuit. In addition, the grid acts as an electrostatic shield so that very little plate energy is fed back to the input circuit through the plate-cathode interelectrode capacitance. Thus, the tendency for the circuit to oscillate is greatly reduced. This makes it possible to use triodes as h-f amplifiers without the danger of oscillation.

Another advantage of the grounded-grid amplifier is the elimination of noise signals. In tetrodes and pentodes the conduction current divides between the screen grid and plate. The division of current is not steady, and random variations occur to cause corresponding plate current changes. These variations are amplified and appear in the amplifier output as noise signals.

With no screen grid, no division of plate current occurs in a triode, and thus this noise is not generated. However, compared to tetrodes, the grounded-grid amplifier has low gain. Therefore, its use is confined mainly to the ultra high frequency applications or where elimination of noise is imperative.

SUPPRESSION OF PARASITIC OSCILLATION

The frequency at which an h-f amplifier operates normally depends upon the inductance and capacitance values of its tank circuits. However, the interelectrode capacitances of the tube circuit may produce resonance in conjunction with the grid and plate circuit leads. These wires can be thought of as one-turn inductors which produce oscillations at frequencies very much higher than the resonant frequency of the tank circuit. Known as **PARASITIC OSCILLATIONS**, these oscillations are very undesirable because they absorb power and thus reduce the available power output of the amplifier.

Referring to Figure 15 as an example, at very high frequencies capacitor C_4 and C_2 offer very low reactance. This may permit the inductance of the connecting wires from the plate to C_2 , C_2 to C_4 , and C_4 to ground to form a tank circuit with the plate-filament capacitance of the tube.

In like manner, the grid circuit may resonate at a frequency determined by the grid-to-filament capacitance and the inductance of the connecting wires from the grid to C_1 , C_1 to C_3 , and C_3 to ground. Thus the complete circuit

forms a tuned-plate tuned-grid oscillator which generates a frequency many times the resonant frequency of the tank circuits.

To prevent parasitic oscillations it is also common practice to increase the distributed inductance of the plate circuit so that it resonates at a frequency lower than that of the grid circuit.

Often this is done by increasing the length of the plate circuit connecting wires or leads. In many cases the extra length is wound into a coil called a PARASITIC CHOKE. At the much lower operating frequency of the amplifier circuit, the reactance of this choke is too small to cause any appreciable effect.

Another system of suppressing parasitic oscillations is to insert small-value resistors in the plate or grid leads. These resistors are inserted directly into the parasitic oscillating circuit, hence they lower the circuit Q so that oscillations cannot be sustained. On the other hand, the resistors are not in the operating frequency tank circuit, but rather in the lead going to it. Therefore, they have no appreciable effect at the desired operating frequency of the amplifier.

In the circuit of Figure 15, such a resistor would be inserted between the plate and point 2, or between the grid and point 1. In either case the resistor would be directly in series with one of the parasitic circuits traced previously. Hence, with a lower Q these circuits would not sustain the parasitic oscillations.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

H-F POWER AMPLIFIERS

GRID DRIVING POWER

PUSH-PULL AMPLIFIERS

FREQUENCY MULTIPLIERS

NEUTRALIZATION

PLATE NEUTRALIZATION

GRID NEUTRALIZATION

THE GROUNDED-GRID AMPLIFIER

SUPPRESSION OF PARASITIC OSCILLATION

33. Does it take signal power to drive a Class B or Class C amplifier?
34. What are some of the advantages of a push-pull amplifier?
35. How does a frequency multiplier differ from a conventional h-f amplifier?
36. What changes are necessary in the circuit of Figure 17 to operate it as a quadrupler?
37. What basic factor in a tube can cause oscillation in an amplifier circuit?
38. Basically how does a neutralization circuit eliminate the effects of grid-to-plate capacity?
39. Give a brief explanation of how the neutralization circuit of Figure 18 operates.
40. In the circuit of Figure 19, C_N would have a value (greater) (equal to) (less) than the value of C_{gp} .
41. Does the Rice or Hazeltine neutralization method require a center tapped plate tank coil?
42. Does a grounded-grid amplifier require neutralization?
43. Compared to a grounded-cathode amplifier, the grounded-grid amplifier has a higher gain. True or False?
44. What are parasitic oscillations?
45. Why are low value resistors often connected in the plate and grid leads of an h-f amplifier?

SUMMARY

Vacuum tube high-frequency amplifiers amplify signal frequencies in the range from 10 kHz to about 1000 MHz. Most high-frequency amplifiers are of the tuned type. Tuning is provided by resonant circuits in the coupling network. Tuned amplifiers may use either capacitive or inductive coupling.

When the coupling circuit employs a single resonant circuit, the method is called single-tuned coupling. This provides a relatively narrow bandpass. In amplifiers that require a wide bandpass, two resonant circuits are employed. This method is called double-tuned coupling. A high-Q resonant circuit provides a large gain while a low-Q circuit results in small gain.

Tetrodes and pentodes are normally used in high-frequency voltage amplifiers due to their high amplification factor and low interelectrode capacity. Where large power is required, triodes are used. Most power amplifiers employ Class C operation. Where distortion cannot be tolerated, Class B operation is used. Many high-frequency amplifiers are operated push-pull to obtain greater power output.

A high-frequency amplifier may be used as a frequency multiplier by tuning the plate circuit to a harmonic of the grid circuit frequency. In this case the output signal is some multiple of the input frequency. Neutralization is employed in high-frequency amplifiers to prevent undesired oscillations due to regenerative feedback from the output to the input circuit. Two methods are in common use. The Hazeltine or plate method and the Rice or grid method. The grounded-grid amplifier eliminates the need for neutralization, but its gain is low.

ESSENTIAL SYMBOLS AND EQUATIONS

A_v = Voltage amplification

E_o = Output signal voltage (volts)

E_{IN} = Input signal voltage (volts)

f = Frequency (hertz)

g_m = Tube transconductance (mhos)

% Eff. = Percent Efficiency

L = Inductance (henries)

P_{IN} = Input power (dc) (watts)

P_o = Output power (signal) (watts)

R = Effective resistance (ohms)

Z_L = Total load impedance (ohms)

λ = Wavelength (meters)

$$\lambda = \frac{300,000,000}{f} \quad (1)$$

$$\lambda = \frac{300}{f(\text{MHz})} \quad (1a)$$

$$f(\text{MHz}) = \frac{300}{\lambda} \quad (1b)$$

$$A_v = \frac{E_o}{E_{IN}} \quad (2)$$

$$A_v = g_m Z_L \quad (3)$$

$$\% \text{ Eff.} = \frac{P_o}{P_{IN}} \times 100 \quad (4)$$

IMPORTANT DEFINITIONS

BANDPASS — All of the frequencies in a tuned circuit which have amplitudes within 70.7% of the peak. Also called the passband.

DOUBLE-TUNED AMPLIFIER — An amplifier in which the coupling network has two tuned circuits.

FREQUENCY MULTIPLIER — A class C amplifier that is operated with its plate circuit tuned to a harmonic of the grid circuit frequency.

GROUNDED-GRID AMPLIFIER — A high-frequency amplifier arrangement that permits the use of a triode tube, without the necessity of neutralization, by grounding the grid instead of the cathode.

IMPEDANCE COUPLING — A capacitive coupling method in which one (or both) of the shunt components is an inductor or an LC circuit.

LINK COUPLING — An h-f coupling system that employs a link circuit composed of a transmission line with terminating loops of one or two turns to transfer the h-f energy from one circuit to another.

MUTUAL IMPEDANCE COUPLING — An inductive coupling method in which the primary and secondary coils are not coupled magnetically. A third coil which is part of both circuits provides the coupling.

NEUTRALIZATION — The process of nullifying the voltage fed back through the interelectrode capacitance of an amplifier tube by providing an equal voltage of opposite polarity.

PARASITIC OSCILLATIONS — A very high frequency generated by interelectrode and stray capacitances in resonance with the lead inductance of the circuit.

PLATE CIRCUIT EFFICIENCY — The ratio of the power developed in the load to the direct current power delivered to the circuit by the dc power supply.

RADIO FREQUENCIES — The band of frequencies extending from 1×10^4 to 3×10^{12} Hz.

RESISTANCE-CAPACITANCE COUPLING — A capacitive coupling method in which both the input and output shunt components are resistors. Also called RESISTANCE COUPLING.

RESISTANCE COUPLING — See RESISTANCE-CAPACITANCE COUPLING.

IMPORTANT DEFINITIONS (Continued)

SINGLE-TUNED AMPLIFIER — An amplifier in which the coupling network has one tuned circuit.

TRANSFORMER COUPLING — An inductive coupling method in which the energy transfer is the result of primary flux cutting the secondary of a transformer.

PRACTICE EXERCISE SOLUTIONS

1. inversely — Wavelength is inversely proportional to frequency as shown

by Equation 1a — $\lambda = \frac{300}{f(\text{MHz})}$.

2. $\lambda = 20$ meters. The wavelength is determined using Equation 1a —

$$\lambda = \frac{300}{f(\text{MHz})} = \frac{300}{15} = 20 \text{ meters.}$$

3. A stage is defined as a group of components and circuits associated with one or more electron tubes.
4. A tuned amplifier is an amplifier with a restricted frequency response that provides maximum gain over a narrow frequency band.
5. Resistance coupling occurs when both input and output shunt components are resistors. If one or both of the shunt components is an inductor or LC circuit, it is impedance coupling.
6. With transformer coupling, the energy transfer is a result of the magnetic flux from the primary winding cutting the secondary. In mutual impedance coupling, the primary and secondary coils are not coupled magnetically, but a third coil which is part of both circuits provides the coupling.
7. high pass — The RC coupling circuit is basically a high pass filter circuit.
8. False — For proper operation the coupling circuit should have a long time constant.
9. C_2 is a bypass capacitor which prevents signal variations from reaching the power supply.
10. The input and output capacities affect the high-frequency response of the circuit. Remember; their reactance decreases as signal frequency increases.
11. Lowering the value of R_2 extends the high frequency response but lowers the overall gain.
12. At the resonant frequency the load impedance increases causing an increase in the output voltage.
13. The lower dc resistance reduces the IR drop across the load. This makes it unnecessary for the plate supply voltage to be much higher than the required plate voltage.

PRACTICE EXERCISE SOLUTIONS (Continued)

14. Yes, the voltage gain of a pentode amplifier is $A_v = g_m Z_L$ while that for

a triode amplifier is $A_v = \frac{\mu R_L}{R_L + r_p}$.

15. The circuit of Figure 7 is a single-tuned amplifier since the coupling circuit contains only one tuned circuit.
16. The Q of the tuned circuit determines the bandpass of the amplifier. With a high Q circuit, the bandpass is narrow and with a low Q circuit the bandpass is wide.
17. The impedance of the tuned circuit and likewise the amplifier gain is proportional to the L to C ratio. With a large L to C ratio, the impedance is high and likewise the gain is high.
18. 70.7% — The bandpass of an amplifier includes the frequencies that are amplified an amount equal to or greater than 70.7% of maximum amplification.
19. False — The transformers used in h-f amplifiers generally employ powdered iron or air cores.
20. When the coupling is increased beyond the critical point, the response curve takes on a double-humped appearance and the bandwidth increases.
21. A steep sided response curve is said to have high skirt selectivity.
22. In a mutual impedance coupling arrangement there is no magnetic coupling between the plate and grid circuit coils. Coupling is accomplished through a third inductor which is common to both circuits.
23. L_3 is the mutual impedance in the circuit of Figure 12.
24. Link coupling provides a means of inductively coupling circuits that are physically separated.
25. The lower interelectrode capacity of a tetrode makes it less subject to oscillation.
26. True — The i-f and r-f amplifiers in a television receiver are voltage amplifiers since no power is required.
27. The efficiency of an amplifier is the ratio of the signal power output to the dc power input from the power supply.

PRACTICE EXERCISE SOLUTIONS (Continued)

28. 53% — The efficiency of an amplifier is equal to:

$$\%E_{\text{eff}} = \frac{P_o}{P_{\text{IN}}} \times 100 = \frac{200}{750 \times .5} \times 100 = \frac{200}{375} \times 100 = 53\%$$

(approx.)

29. The pulses of plate current develop a circulating current in the tank circuit which is in the form of a sinewave.
30. (low) — To obtain a high efficiency the plate voltage must be low when the plate current pulses occur so the dc power input is low.
31. A Class B amplifier can be operated Class C by increasing the value of grid bias.
32. False — It is common practice to ground the center tap of the filament supply by using a center tapped transformer or a resistor network.
33. Yes, a Class B or Class C amplifier requires grid driving power since the grid must be driven positive to produce grid current.
34. A push-pull amplifier produces twice the power output of a single-ended stage and cancels even harmonic distortion.
35. A frequency multiplier is basically a Class C h-f amplifier. The plate tank circuit, however, is tuned to the desired harmonic of the input signal frequency.
36. To operate the circuit of Figure 17 as a quadrupler, the plate tank circuit would be tuned to the fourth harmonic, 20 MHz, of the input signal.
37. The grid-to-plate capacity of a tube is the factor that causes oscillation.
38. The neutralization circuit eliminates the effect of grid-to-plate capacity by applying a feedback voltage of the same magnitude but opposite phase which “bucks” the signal fed back through the grid-to-plate capacity.
39. In the circuit of Figure 18, C_N feeds a signal back to the grid which is 180° out of phase with the signal fed back through C_{gp} . The 180° out of phase signal is obtained by center tapping the plate tank circuit.
40. equal to — The value of the neutralization capacitor should be approximately equal to the value of the grid-to-plate capacity.

PRACTICE EXERCISE SOLUTIONS (Continued)

- 41. The Hazeltine or plate neutralization method employs a center tapped plate tank coil.**
- 42. No, in most cases a grounded-grid amplifier can be operated as a h-f amplifier without the need for neutralization.**
- 43. False — A grounded-grid amplifier has a lower gain than a grounded-cathode amplifier.**
- 44. Parasitic oscillations are oscillations that occur at frequencies other than the resonant frequency of the tuned circuits.**
- 45. The low value resistors lower the Q of the grid and plate circuit loads, thus preventing parasitic oscillation.**

NOTICE: To keep our lessons up-to-date, we add material to the texts from time to time, and make occasional changes in the examination questions. Therefore, when checking your answers with this check sheet, look at the code LETTER (A, B, C, etc.) on your examination sheet to determine which of the following groups of answers apply to the lesson which you have.

1088C

1. B - WHAT IS THE FREQUENCY OF A SIGNAL THAT HAS A WAVELENGTH OF 25 METERS -- 12 MHz.
The frequency is found by rearranging the wavelength formula.

$$\lambda = \frac{300}{f(\text{MHz})} \quad f(\text{MHz}) = \frac{300}{\lambda} = \frac{300}{25} = 12 \text{ MHz}$$

2. B - TO COUPLE MAXIMUM ENERGY FROM THE V_1 STAGE TO THE V_2 STAGE IN THE CIRCUIT OF FIGURE 3 -- THE C_3R_3 TIME CONSTANT SHOULD BE LONG.

The coupling circuit time constant should be long over the desired frequency range.

3. C - MAXIMUM GAIN OCCURS IN THE CIRCUIT OF FIGURE 7 -- AT THE RESONANT FREQUENCY OF THE OUTPUT CIRCUIT ($C_3L_1C_0C_{IN}$).

Maximum gain occurs when the impedance of the output circuit is highest, which is at the resonant frequency of C_3 , L_1 and the shunt capacitances of C_0 and C_{IN} .

4. D - WHAT IS THE GAIN OF THE CIRCUIT OF FIGURE 7 IF V_1 HAS A g_m OF 5,000 μmhos AND THE TANK CIRCUIT IMPEDANCE AT RESONANCE IS 47,000 OHMS -- 235.

The stage gain is equal to

$$A_V = g_m Z_L = 5,000 \times 10^{-6} \times 4.7 \times 10^4 = 5 \times 10^{-3} \times 4.7 \times 10^4 = 23.5 \times 10^1 = 235$$

5. A - WHEN THE TANK CIRCUIT IN A TUNED AMPLIFIER HAS A HIGH Q -- THE GAIN IS HIGH AND THE BANDPASS IS NARROW.

A high Q results in a high load impedance which yields a high gain and a narrow bandpass.

6. A - IN A DOUBLE-TUNED INDUCTIVELY COUPLED AMPLIFIER WHEN THE COUPLING IS INCREASED BEYOND THE CRITICAL POINT, -- THE BANDWIDTH INCREASES.

Increasing the coupling beyond the critical point produces a double humped response curve which increases the bandwidth.

7. A - WHAT IS THE EFFICIENCY OF A POWER AMPLIFIER IF THE OUTPUT POWER IS 50 WATTS, $E_b = 700$ VOLTS AND $I_b = 120 \text{ mA}$? -- 60%.

The efficiency is found as follows:

$$P_{IN} = E_b \times I_b = 7 \times 10^2 \times 1.2 \times 10^{-1} = 84 \text{ watts}$$

$$E_{ff} = \frac{P_O}{P_{IN}} \times 100 = \frac{50}{84} \times 100 = 60\% \text{ (approx.)}$$

8. B - A FREQUENCY MULTIPLIER STAGE -- OPERATES CLASS C.

To obtain the high harmonic content in the plate circuit, a frequency multiplier is operated as Class C.

9. B - IN THE CIRCUIT OF FIGURE 18 -- THE SIGNALS COUPLED THROUGH C_{gp} AND C_N ARE 180° OUT OF PHASE.

The signal coupled through C_N is the same magnitude but 180° out of phase with the signal coupled through C_{gp} .

10. A - THE CIRCUIT OF FIGURE 20 -- REQUIRES NO NEUTRALIZATION.

Due to the shielding effect of the grid, the grounded-grid amplifier requires no neutralization.

1088B

All explanations are the same as for 1088C except for that given below.

3. C - MAXIMUM GAIN OCCURS IN THE CIRCUIT OF FIGURE 7 -- AT THE RESONANT FREQUENCY OF THE OUTPUT CIRCUIT.

Maximum gain occurs when the impedance of the C_3L_1 tank circuit is highest, which is at the resonant frequency of C_3 and L_1 .

PRACTICE EXERCISE SOLUTIONS (Continued)

- 41. The Hazeltine or plate neutralization method employs a center tapped plate tank coil.**
- 42. No, in most cases a grounded-grid amplifier can be operated as a h-f amplifier without the need for neutralization.**
- 43. False — A grounded-grid amplifier has a lower gain than a grounded-cathode amplifier.**
- 44. Parasitic oscillations are oscillations that occur at frequencies other than the resonant frequency of the tuned circuits.**
- 45. The low value resistors lower the Q of the grid and plate circuit loads, thus preventing parasitic oscillation.**

1088A

All explanations are the same as for 1088C except for that given below.

3. C - MAXIMUM GAIN OCCURS IN THE CIRCUIT OF FIGURE 7 -- AT THE RESONANT FREQUENCY OF C_3 AND L_1 .
Maximum gain occurs when the impedance of the C_3L_1 tank circuit is highest, which is at the resonant frequency of C_3 and L_1 .

PRACTICE EXERCISE SOLUTIONS (Continued)

- 41. The Hazeltine or plate neutralization method employs a center tapped plate tank coil.**
- 42. No, in most cases a grounded-grid amplifier can be operated as a h-f amplifier without the need for neutralization.**
- 43. False — A grounded-grid amplifier has a lower gain than a grounded-cathode amplifier.**
- 44. Parasitic oscillations are oscillations that occur at frequencies other than the resonant frequency of the tuned circuits.**
- 45. The low value resistors lower the Q of the grid and plate circuit loads, thus preventing parasitic oscillation.**

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1088C

Example: A library lends

A	☒
B	☐
C	☐
D	☐

A. books. B. automobiles. C. boots. D. airplanes.

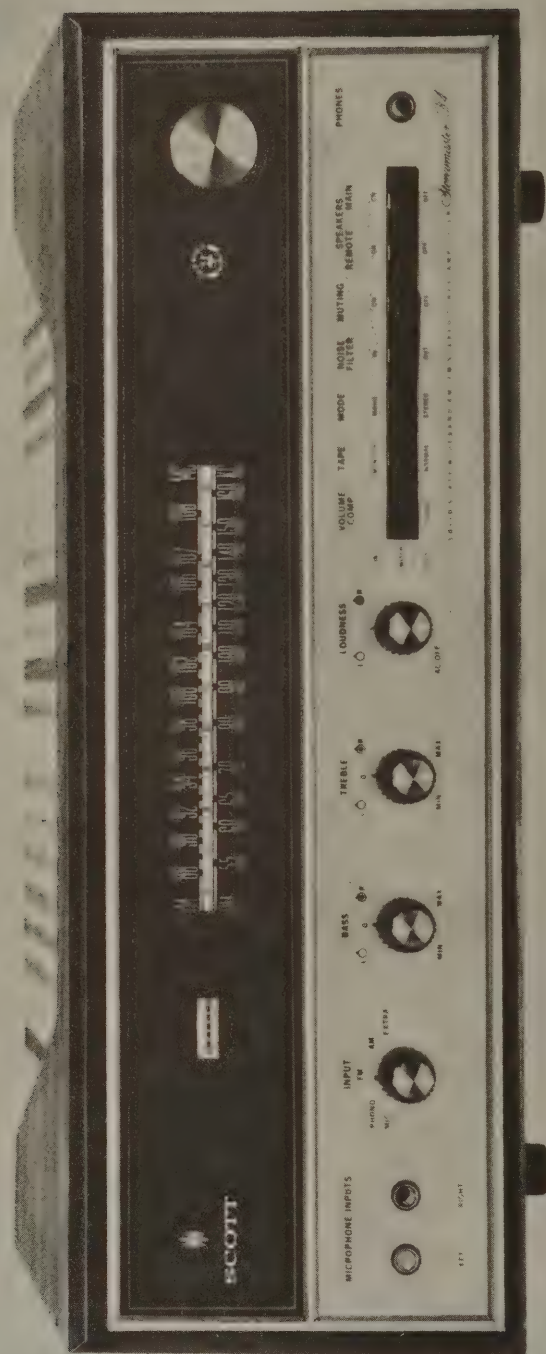
1. A ☐ B ☒ C ☐ D ☐ What is the frequency of a signal that has a wavelength of 25 meters?
(A) 24 MHz. (B) 12 MHz. (C) 6 MHz. (D) 3 MHz.
2. A ☐ B ☒ C ☐ D ☐ To couple maximum energy from the V_1 stage to the V_2 stage in the circuit of Figure 3
(A) the C_3R_3 time constant should be short. (B) the C_3R_3 time constant should be long. (C) the value of R_3 should be very small. (D) the value of C_2 must be very small.
3. A ☐ B ☐ C ☒ D ☐ Maximum gain occurs in the circuit of Figure 7
(A) below the resonant frequency of C_3 and L_1 . (B) above the resonant frequency of C_3 and L_1 . (C) at the resonant frequency of the output circuit ($C_3L_1C_0C_{IN}$). (D) at all frequencies.
4. A ☐ B ☐ C ☐ D ☒ What is the gain of the circuit of Figure 7 if V_1 has a g_m of 5,000 μ mhos and the tank circuit impedance at resonance is 47,000 ohms
(A) 57,000. (B) 1,600. (C) 23.5. (D) 235.
5. A ☒ B ☐ C ☐ D ☐ When the tank circuit in a tuned amplifier has a high Q
(A) the gain is high and the bandpass is narrow. (B) the gain is low and the bandpass is wide. (C) the gain is low and the bandpass is narrow. (D) the gain is high and the bandpass is wide.
6. A ☒ B ☐ C ☐ D ☐ In a double-tuned inductively coupled amplifier when the coupling is increased beyond the critical point,
(A) the bandwidth increases. (B) the bandwidth decreases. (C) the gain increases. (D) the bandwidth remains the same.
7. A ☒ B ☐ C ☐ D ☐ What is the efficiency of a power amplifier if the output power is 50 watts, $E_b = 700$ volts and $I_b = 120$ mA?
(A) 60%. (B) 38%. (C) 50%. (D) 90%.
8. A ☐ B ☒ C ☐ D ☐ A frequency multiplier stage
(A) generally operates Class A. (B) operates Class C. (C) has a high efficiency. (D) has both the input and output circuits tuned to the same frequency.
9. A ☐ B ☒ C ☐ D ☐ In the circuit of Figure 18
(A) the signals coupled through C_{gp} and C_N are in phase. (B) the signals coupled through C_{gp} and C_N are 180° out of phase. (C) the value of C_N is usually many times greater than C_{gp} . (D) the signals at point P and N are in phase.
10. A ☒ B ☐ C ☐ D ☐ The circuit of Figure 20
(A) requires no neutralization. (B) has a very high gain. (C) requires neutralization. (D) is a grounded cathode amplifier.

TRANSISTOR PUSH-PULL AND HIGH FREQUENCY AMPLIFIERS

1091

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





The AM/FM stereo receiver shown above contains a wide variety of amplifier types. The audio output stages are generally Class AB push-pull amplifiers while the audio preamplifier stages are generally single-ended Class A stages. The i-f and r-f amplifiers are likewise operated Class A for minimum distortion.

Courtesy H. H. Scott

CONTENTS

Push-pull Transistor Amplifiers	Page 4
Class B Push-pull Operation	Page 4
Class A Push-pull Operation	Page 6
Class AB Push-pull Operation	Page 9
Phase Inverters	Page 10
Limitations of RC Coupling	Page 13
Complementary-Symmetry Circuit	Page 13
Direct Coupled Circuit	Page 15
High-frequency Amplifiers	Page 15
Tuned Amplifiers	Page 16
Neutralization	Page 24
Push-pull and Parallel H-F Amplifiers	Page 26
High-frequency Components	Page 28
Impedance Matching	Page 29
Output Transformer Coupling	Page 30
Tuned Output Coupling	Page 31

*On how to throw a curve; "You can only do it when you pitch,
and pitch and pitch."*

—Christy Mathewson

TRANSISTOR PUSH-PULL AND HIGH FREQUENCY AMPLIFIERS

Transistors are used for many purposes. Some are designed for high power operation, some for high frequency operation, and so on. A given transistor will do its job better if the circuit in which it is used is designed especially for that job. In this lesson you will study amplifier circuits which are designed for specific purposes, such as for amplifying a certain band of frequencies. You will see how the use of a tank circuit makes this possible.

PUSH-PULL TRANSISTOR AMPLIFIERS

A push-pull circuit is an amplifier circuit using two transistors connected so that each contributes output power to a load. The advantages of push-pull operation are many. Among the most important are increased power output, reduction in even harmonic distortion and noise level, and greater efficiency. Push-pull operation is used with Class A, Class B and Class AB stages to reduce distortion.

MOST TRANSISTOR PUSH-PULL AMPLIFIERS REQUIRE INPUT SIGNALS WHICH ARE 180° OUT OF PHASE WITH EACH OTHER. That is, when the signal to one transistor is increasing in the negative direction, the signal to the other transistor is increasing in the positive direction. Circuits which provide this 180° phase shift are described later in the lesson.

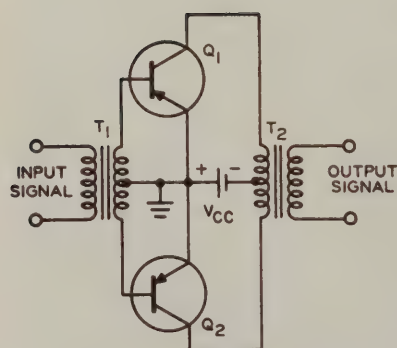


Figure 1

Class B Push-pull Operation

The circuit of a Class B push-pull amplifier is shown in Figure 1. This type of push-pull amplifier requires that the input signals to transistors Q1 and Q2 be 180° out-of-phase with each other. The secondary of input transformer T1 is center tapped to provide the 180° phase reversal.

Proper operating voltages are applied to the collectors of Q1 and Q2 by connecting the center tap of output transformer T2 to the negative terminal of VCC. The dc resistance of T2 is very small so that the voltage at the collectors of Q1 and Q2 with no signal applied is essentially the same as that at the negative terminal of VCC.

With no signal, the bases of Q1 and Q2 in this circuit are operated at zero base bias to provide Class B operation. The center tap of the T1 secondary (

is connected to a ground point, as are both emitters. Thus the no-signal voltage at each base is the same as at the emitter, and the base-emitter bias is zero.

Figure 2 shows the transfer characteristic curves for the transistor circuit of Figure 1. Complete collector current cutoff is assumed at zero base current. The curve for Q_2 has been turned upside down and placed below the Q_1 curve. The resultant $I_C - I_B$ transfer characteristic curve is the line formed by joining the two curves as shown in Figure 2. **NOTICE THAT THE ZERO BASE CURRENTS ON BOTH CURVES ARE ALIGNED.**

The signal alternation that produces base current in Q_1 keeps Q_2 cut off. The collector current pulse of Q_1 has the waveshape shown as i_{c1} . The next alternation causes Q_2 to conduct, as shown by the lower alternation labeled i_{c2} . The combination of i_{c1} and i_{c2} produces a waveshape that is distorted due to the bends in the $I_C - I_B$ curves. This type of distortion is called **CROSSOVER DISTORTION**. It occurs at the point where collector current crosses over from one transistor to the other. The output signal induced at the secondary of T_2 is a replica of this wave made up of i_{c1} and i_{c2} currents.

The crossover distortion shown in Figure 2 may be reduced by applying a small amount of forward bias to the transistors. This places the operating point a little above cutoff. The resultant of the two curves then gives a more linear operation near cutoff.

Figure 3 shows a Class B push-pull amplifier with a small forward bias applied to the base-emitter circuits of the two transistors. The voltage divider composed of R_1 and R_2 makes point X negative with respect to the emitters. The value of R_2 is large compared to the value of R_1 . Hence the forward bias voltage developed across resistor R_1 is very small and the resulting no-signal base currents are small. Since the left end of R_1 connects to the center tap of the secondary of T_1 , base currents flow through the secondary of T_1 . R_2 carries the base currents of both transistors as well as the current which flows through R_1 .

In this base bias network, electrons flow from the negative side of V_{CC} through R_2 to point X. Part of this current flows through R_1 back to the positive terminal of V_{CC} . Part of the remainder flows from point X through the upper half of the secondary of T_1 to the base of Q_1 , and the rest is through the lower half of the secondary of T_1 to the base of Q_2 . Electrons leave each emitter to return to the positive terminal of battery V_{CC} , which is connected to ground. The values of R_2 and R_1 are chosen so the small voltage at point X remains steady in normal operation.

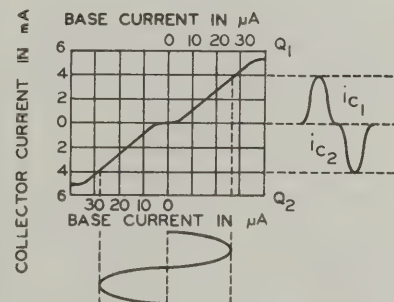


Figure 2

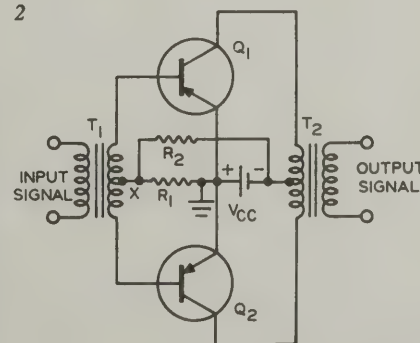


Figure 3

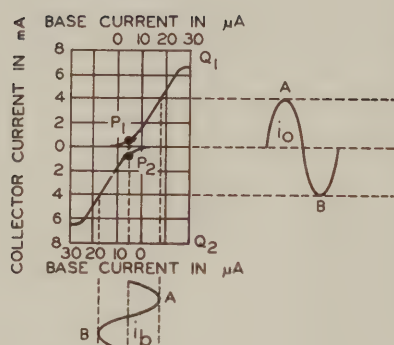


Figure 4

The effect of the small bias is easily illustrated by examining the transfer characteristics of this circuit as shown in Figure 4. The curve of Q_1 is drawn at the top and the Q_2 curve is at the bottom. The small amount of forward bias places the operating point for Q_1 at P_1 and for Q_2 at P_2 . Comparing Figure 4 with Figure 2 shows that both curves are shifted due to the addition of forward bias. The zero base current points are no longer aligned. **THE NO-SIGNAL BASE CURRENT POINTS ARE NOW ONE UNDER THE OTHER.** Both transistors are operated with a no-signal base current or bias of $5 \mu A$.

Notice that the two curves of Figure 4 are connected by a dashed line. This represents the net current in the transformer. Because the collector currents in the transformer are out-of-phase, only the **DIFFERENCE** between the two currents produces an output. Thus, when both collector currents are the same, the fields in the transformer cancel each other. The result is the same as zero current at the operating point. The dashed line then indicates the combined operation curve.

On the alternation of the input signal i_b labeled A, the collector current of Q_1 increases to 4 mA, while Q_2 is cut off. On the alternation of i_b labeled B, Q_1 is cut off and the collector current of Q_2 increases to 4 mA. If each point along the combined curve, including the dashed line section, is used to plot the resulting current in the transformer, the i_o wave of Figure 4 results. The crossover distortion no longer can be seen. The output signal from T_2 of Figure 3 will be a replica of this i_o curve. The result of providing a small forward bias is a reduction of distortion.

The curve running from the lower left to the upper right in the graph of Figure 4, including the dashed line portion, is called the **COMPOSITE TRANSFER CHARACTERISTIC CURVE**. This curve represents the combined characteristics of Q_1 and Q_2 . The composite curve is plotted by adding the individual Q_1 and Q_2 curves together, considering one as positive and the other as negative. The dashed lines running from the composite curve to the right are used to plot the output waveshape i_o . The i_o waveshape is a result of the combined collector currents of Q_1 and Q_2 . The curve of Figure 2 is also a composite transfer characteristic curve.

The composite characteristic curves for push-pull operation are formed by turning one curve upside down and then aligning the operating points so that the operating point on one curve is directly beneath the operating point on the other. The two curves are then connected as explained above.

Class A Push-pull Operation

When transistors are operated in Class B push-pull circuits, the distortion of single-ended Class B circuits is greatly reduced. By applying a small forward

base bias even the crossover distortion can be made small. By further increasing the forward base bias until Class A operation is reached, distortion is reduced to a minimum. The power output is about twice that of a single-ended Class A stage.

Figure 5 shows the circuit of a typical Class A push-pull amplifier. This circuit includes the bypass capacitors and temperature stabilizing resistors normally found in a practical circuit. Transistors Q_1 and Q_2 are PNP transistors. To forward bias the base-emitter junctions, the bases are made negative with respect to the emitters. The emitters are connected through thermal runaway circuits R_1C_1 and R_2C_2 to the positive terminal of V_{CC} . The values of R_1 and R_2 are chosen so that at normal operating currents the voltages developed across them are very small. Thus, the emitters of both transistors are at approximately the potential of the positive terminal of V_{CC} .

The bases of Q_1 and Q_2 are connected to opposite ends of transformer T_1 . The dc resistance of the secondary of T_1 is very small. With no signal applied, the voltage at the two bases is that of the center tap of the T_1 secondary.

Resistors R_3 and R_4 of Figure 5 form a voltage divider across V_{CC} . One end of this divider connects to the negative terminal of V_{CC} , and the other end is connected to ground. Notice that the positive terminal of V_{CC} is also grounded. The center tap of T_1 is connected to the point between voltage divider resistors R_3 and R_4 . This point is negative with respect to the positive terminal of V_{CC} . Thus, the base-emitter junctions of the two transistors are forward biased. Capacitor C_3 provides a path to ground for ac signal currents but blocks dc so that the base bias is not shorted to ground.

Resistors R_1 and R_2 and capacitors C_1 , C_2 and C_3 could also be used in the Class B circuit of Figure 3. They were omitted to simplify the explanation.

Collector currents for Q_1 and Q_2 of Figure 5 follow similar paths. For Q_1 , collector current flows from the negative side of V_{CC} through the upper half of T_2 to the collector of Q_1 . Current leaves Q_1 at the emitter and flows through the R_1C_1 thermal runaway circuit to the positive terminal of V_{CC} .

In the Q_2 circuit electrons flow from the negative side of V_{CC} , through the lower half of T_2 to the collector of Q_2 . Electrons leave Q_2 at the emitter and flow through the R_2C_2 thermal runaway circuit to the positive terminal of V_{CC} .

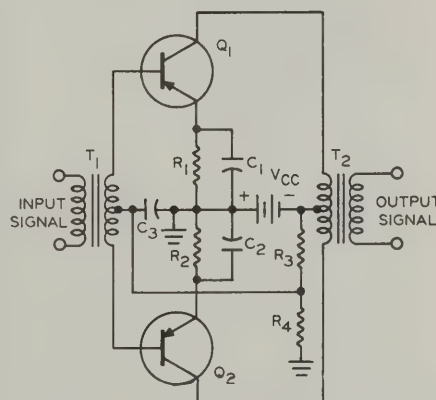


Figure 5

When the input signal applied to the primary of T_1 induces a voltage in the upper half of the T_1 secondary which aids the forward bias on Q_1 , the collector current of Q_1 increases. When the input signal induces a voltage in the upper half of the T_1 secondary which opposes forward bias of Q_1 , the collector current of Q_1 decreases.

When the input signal applied to the primary of T_1 induces a voltage in the secondary which aids the forward bias on Q_2 , the collector current of Q_2 increases. When the input signal induces a voltage in the lower half of the T_1 secondary which opposes the forward bias on Q_2 , the collector current of Q_2 decreases.

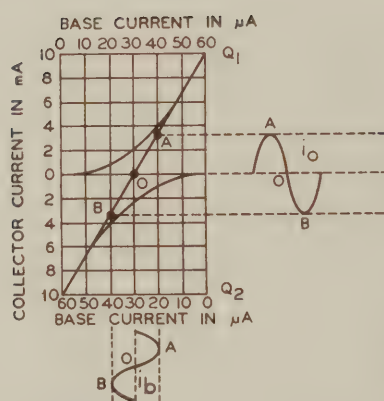
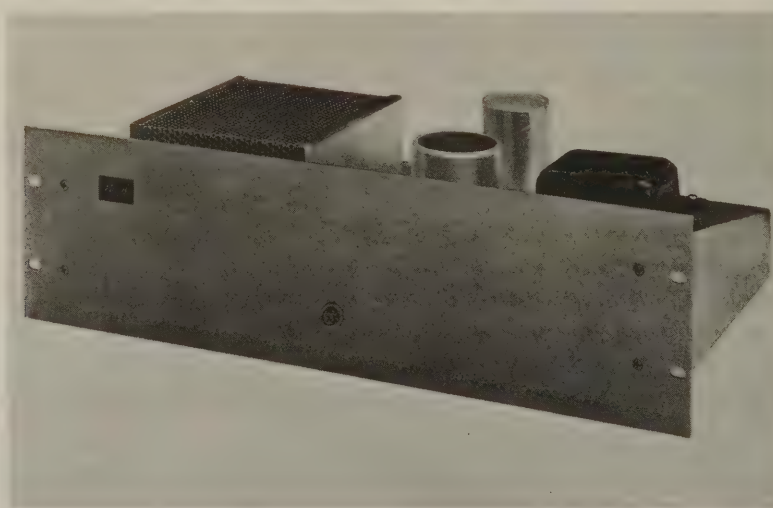


Figure 6

The transfer characteristic curves for the circuit of Figure 5 are shown in Figure 6. The curve for Q_1 is drawn in the upper half of the graph. The curve for Q_2 is an inverted form of the Q_1 curve, and is drawn in the lower half of the graph.

Figure 6 shows how the Class A transfer curves are combined to obtain a composite curve. Let's assume that the values chosen for R_3 and R_4 in Figure 5 allow a no-signal base bias current of $30 \mu A$ to flow in both transistors. Notice that the $30 \mu A$ base current lines on the two curves are aligned one over the other. If the two curves do not meet, an extension is drawn between the curves in the form of a straight line connecting them. The resulting composite curve is much straighter than either transistor curve by itself.



This audio power amplifier is designed to deliver an output of 50 watts and utilizes a push-pull output stage. An impedance matching arrangement permits a load impedance of 8 or 100 ohms.

Courtesy R. T. Bozak Mfg. Co.

As the collector current of one transistor decreases, the other increases, so the result is the composite curve of Figure 6. This curve is a very straight line and the distortion is very small. The output i_o is the combined current change in the primary of T_2 . The shape of i_o is found by projecting from the i_b curve at the bottom up to the composite curve and then drawing lines to the right. This is shown for points A, B and O.

In practice, transfer curves are usually much straighter than those shown for Q_1 and Q_2 of Figure 6. These curves are exaggerated to illustrate the effect of the composite curve.

The currents obtained from a composite transfer characteristic curve are not the individual currents flowing in the transistors. The composite current values represent the difference between the currents of the two transistors.

The voltage induced in the secondary of the output transformer is a result of the net current change in the primary. Because the collector currents in the primary are in opposite directions, the difference between them, at any instant, is the important fact. This is why the signal at one base must be the reverse of the signal at the opposite base.

The collector current through one half of the output transformer is always increasing while the collector current in the other is decreasing. As a result, rather than cancelling each other, both collector currents contribute to the output signal at any instant. Any distortion or noise, however, which is generated by both transistors will be cancelled because it will appear in both halves of the output transformer primary with the same polarity at the same instant. Thus, the net change of primary current as a result of such unwanted action will be zero and will not appear in the secondary.

Class AB Push-pull Operation

In the previous examples, we illustrated Class A and Class B operation. Class AB operation is between these two classes of operation. Class AB uses more forward bias than for Class B operation, but less than for Class A operation. As a result, the transfer curves are shifted part way between the extremes of Figures 2 and 6. In using Class AB operation some of the characteristics of both Class A and B are retained. A fairly large signal can be handled with a reasonably small amount of distortion. Class AB operation is sometimes broken down into subclasses to indicate whether it is closer to Class A or Class B.

Phase Inverters

In most push-pull amplifier circuits, the input signal must drive the bases 180° out of phase with respect to each other. A circuit that will provide this type of input signal to a push-pull stage is called a **PHASE INVERTER**.

The simplest form of phase inverter is the coupling transformer used in Figures 1, 3 and 5. The center tap is located at the electrical center of the secondary winding. Equal voltages are therefore produced in EACH HALF OF THE SECONDARY WINDING. When the voltage induced in the secondary makes the upper end positive with respect to the lower end, the upper end is also positive with respect to the tap. Also, at this time, the tap is positive with respect to the lower end. Considering the tap as a reference point, the voltages at each end are 180° out-of-phase with each other.

Transformers can cause distortion if the core saturates when normal signals are applied. If a circuit is to amplify with little distortion, high quality phase inverter transformers must be used. The transformer must also serve as a load impedance for the preceding stage. However, a transformer's impedance falls off at the lower frequencies because the inductive reactance of the windings decreases. In most cases this results in a decreased signal at the primary. Hence, good low-frequency response is only obtained in transformers with large inductance. This in turn requires a large core with many windings and increases the cost and size of the transformer.

At high frequencies, signal losses and resonant conditions are produced by the DISTRIBUTED CAPACITANCE of the winding. This capacitance presents a low-reactance path which shunts the higher frequency signals around the primary. Flux leakage (that part of the magnetic flux which does not cut all the turns of the secondary winding) also produces losses. This is most noticeable at the higher frequencies. This flux leakage due to incomplete linkages is called the LEAKAGE INDUCTANCE. It has the same effect as a small inductance in series with the transformer.

As the signal frequency increases, the inductive reactance of the leakage inductance increases, and more signal is lost across this leakage inductance. Therefore, for good high-frequency response, a transformer must have both low leakage inductance and low distributed capacitance.

Phase inverter circuits employing resistance-capacitance coupling methods may be used where the size, weight, and cost of an input transformer cannot be tolerated. These RC phase inverters are only useful in circuits that are operated Class A or AB; that is, where the transistor base current is well above cutoff.

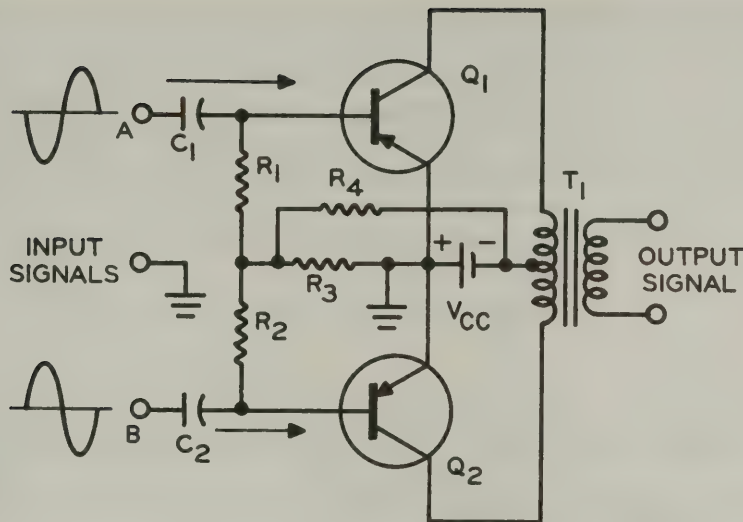


Figure 7

Figure 7 shows a push-pull stage using capacitive input coupling. Compare the input part of this circuit with Figures 1, 3 and 5. In Figure 7, two signals of opposite phase must be available for the input. These signals are produced by a phase inverter.

Figure 8 shows a transistor phase inverter commonly known as a **SPLIT-LOAD PHASE INVERTER** which uses a single input signal to obtain two output signals. This circuit is very similar to a basic common-emitter amplifier. The difference is that the collector load is split into two equal parts. R_2 is one part and R_4 is the other, each providing one-half the total load resistance. The dc voltage across R_4 makes the emitter more negative than in an ordinary amplifier circuit. Therefore, the resistance of R_1 must be larger than usual to make the base more negative than the emitter to obtain the desired forward bias.

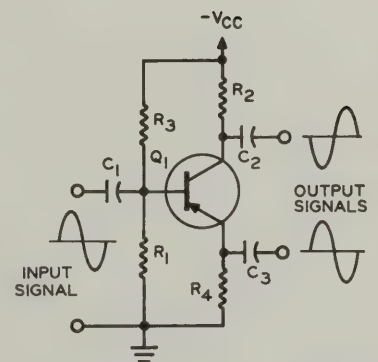


Figure 8

When the input signal swings negative to increase the forward bias, Q_1 collector current increases. With Q_1 collector current increasing, the voltages across R_2 and R_4 increase. Thus, the voltage from the Q_1 collector to ground becomes less negative (more positive) and the voltage from the Q_1 emitter to ground becomes more negative.

When the input signal decreases the forward bias on Q_1 , Q_1 collector current decreases. This causes the voltages across R_2 and R_4 to decrease. Q_1 collector voltage becomes more negative and Q_1 emitter voltage becomes less negative (more positive). Hence, the collector and emitter signal voltages are 180° out-of-phase. The signal voltages coupled by C_2 and C_3 to the circuit output terminals are always 180° out-of-phase with each other.

A two-transistor phase inverter is shown in Figure 9. The Q_1 circuit amplifies and inverts the input signal. Transistor Q_1 is connected as a common-emitter amplifier. A common-emitter amplifier provides a 180° phase reversal between input and output signal voltages. C_5 couples the collector signal to the upper output terminal. A part of this same signal is coupled through R_7 and dc blocking capacitor C_2 to the base of Q_2 . The Q_2 circuit amplifies and inverts this signal, so that the collector signal coupled by C_6 to the lower output terminal is 180° out-of-phase with the signal on the upper output terminal.

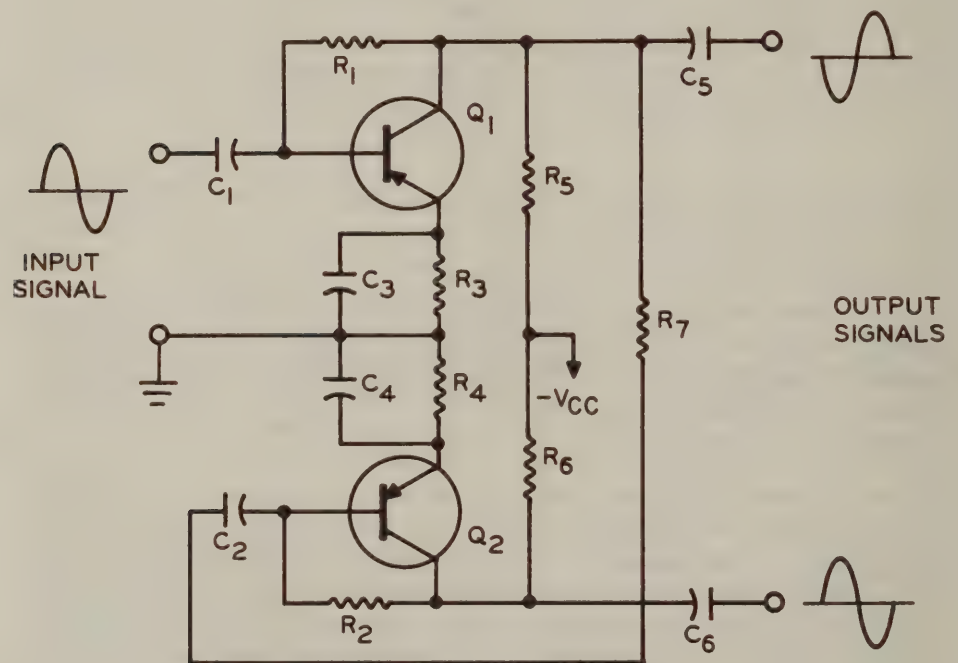


Figure 9

Both transistor circuits in Figure 9 provide the same amount of amplification. Most of the Q_1 collector signal appears across the large series resistor R_7 in the Q_1 -to- Q_2 coupling circuit. The remainder is applied to the Q_2 base and is equal to the signal applied to the base of Q_1 . Therefore, the voltage at the upper output terminal is equal in amplitude to that at the lower output terminal, but 180° out of phase with it.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

PUSH-PULL TRANSISTOR AMPLIFIERS

CLASS B PUSH-PULL OPERATION

CLASS A PUSH-PULL OPERATION

CLASS AB PUSH-PULL OPERATION

PHASE INVERTERS

1. What type of input signals are required for a push-pull amplifier?
2. Why does crossover distortion occur in a Class B push-pull amplifier?
3. How can crossover distortion be eliminated in a push-pull amplifier?
4. Push-pull operation reduces odd harmonic distortion. True or False?
5. R_1 and R_2 in the circuit of Figure 5 provide (forward) (reverse) bias?
6. In a Class A push-pull amplifier, what is the net current in the output transformer primary with no signal?
7. How does the output power of a push-pull stage compare to that of a single-ended stage?
8. What are the factors which affect the low and high-frequency response of a phase inverter transformer?
9. What is the name given to the circuit of Figure 8?
10. When the input signal in the circuit of Figure 8 drives the base of Q_1 more positive, the emitter voltage becomes (more) (less) negative and the collector voltage becomes (more) (less) negative.
11. Would the circuit of Figure 8 have a high voltage gain?
12. In the circuit of Figure 9, what is the phase relation between the input signal and the signal at the base of Q_2 ?
13. Would you expect the circuit of Figure 9 to have a higher gain than the circuit of Figure 8?

Limitations of RC Coupling

Resistance-capacitance coupling cannot always be used to replace transformer coupling. Problems arise when transistors are operated so that base current does not flow at all times. This is typified by Class B operation. In this type of operation the capacitors used with RC coupling develop a bias between base and emitter that changes the operating point when a signal is applied.

The way that this undesirable bias is developed can be illustrated with the push-pull transistor circuit of Figure 7. The input signals applied to this circuit are obtained from a phase inverter. One input signal is applied between terminal A and ground. The other input signal is applied between terminal B and ground. Here, the input signals make the upper input terminal negative with respect to ground at the same time that they make the lower terminal positive with respect to ground.

Assume that Q_1 and Q_2 of Figure 7 are operated in Class B with the R_3 , R_4 forward bias circuit removed. Temporarily we will consider R_3 shorted (to provide a path for the input signal to ground). On the alternation of input signal that forward biases its base-emitter junction, Q_1 conducts base current. With the Q_1 base-emitter junction forward biased, its resistance decreases. Electrons flow in the direction of the arrow and quickly charge C_1 so that the left side is negative with respect to the right side, which is connected to the base.

During the alternation of the input signal that makes the upper input terminal positive, Q_1 does not conduct base current. C_1 is charged and cannot quickly discharge through the high resistance of the reverse biased base-emitter junction. The only discharge path for C_1 is through R_1 in series with the signal source. R_1 is usually a large resistance, so the time constant of the discharge circuit is long and C_1 does not discharge to any extent between the negative alternations. Therefore, a voltage near the peak value of the applied signal is maintained across C_1 . A similar action occurs in the input circuit of Q_2 .

The voltages produced across C_1 and C_2 are in series with the input signal. The capacitor voltages reverse bias the base-emitter junctions of Q_1 and Q_2 except at the peaks of the input signal. This is undesirable because the transistors then operate Class C rather than Class B.

Complementary-Symmetry Circuit

A transistor push-pull circuit which does not require a phase inverter is the **COMPLEMENTARY-SYMMETRY CIRCUIT**. This circuit requires one PNP transistor and one NPN transistor. A typical circuit of this type is

shown in Figure 10. Q_1 is a PNP transistor and Q_2 is an NPN transistor. V_{CC1} supplies negative voltage to the collector of Q_1 , and V_{CC2} supplies the collector of Q_2 with a positive voltage.

Notice that the collector supply voltages are connected series aiding. Resistors R_2 , R_1 , and R_3 form a voltage divider across the two batteries. The resistance of R_2 is usually equal to the resistance of R_3 . Most of the V_{CC1} voltage is across R_2 and most of the V_{CC2} voltage is across R_3 . R_1 has a very low resistance in comparison to R_2 and R_3 , and the voltage across R_1 is very small. R_1 is supplied with a positive voltage at one end and an equal negative voltage at the other, and the voltage at its electrical center is zero or ground potential. Because of this "phantom" ground connection at the center of R_1 , the voltage across the upper half of R_1 is applied between the base and emitter of Q_1 , and the voltage across the lower half of R_1 is applied between the base and emitter of Q_2 . Thus, the voltage divider action of R_1 , R_2 and R_3 provides the proper bias for Q_1 and Q_2 .

Capacitor C_1 couples the input signal to the base of Q_1 and through R_1 to the base of each transistor since the value of R_1 is small. The signal is applied with the same phase to both bases.

The alternations of input signal which reverse bias the base-emitter junction of Q_1 forward bias the base-emitter junction of Q_2 . In Class B or Class AB operation, this polarity of input signal drives Q_1 into cutoff and Q_2 into conduction. Q_2 collector current consists of electron flow from the negative terminal of V_{CC2} through the transformer T_1 primary to the emitter of Q_2 . Electrons flow from emitter to the collector of Q_2 , to the positive terminal of V_{CC2} .

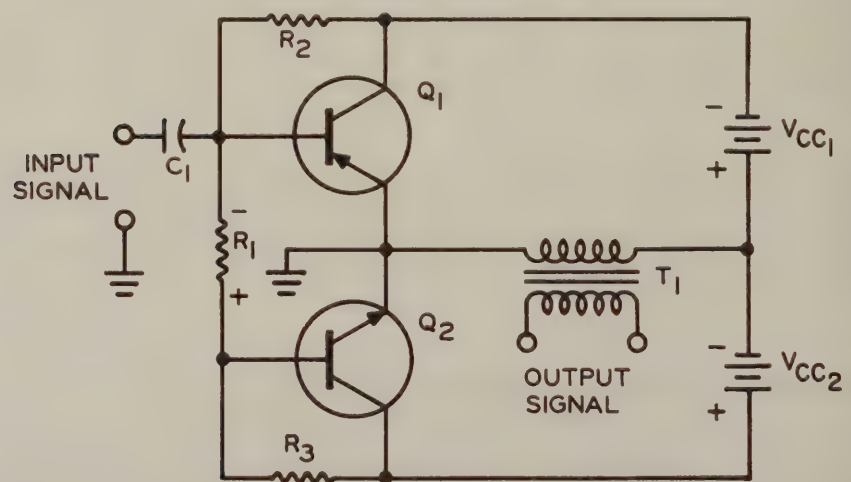


Figure 10

The opposite alternations of input signal forward bias Q_1 and reverse bias Q_2 . Thus, Q_1 conducts and Q_2 is driven into cutoff. Q_1 collector current consists of electron flow from the negative terminal of V_{CC1} to the collector of Q_1 . The path is completed by electron flow from the collector to the emitter of Q_1 , through the primary of transformer T_1 and back to the positive side of V_{CC1} . Thus, the output signal is developed by alternate conduction of transistors Q_1 and Q_2 through the primary of the output transformer.

Notice that the current in the T_1 primary reverses on each half cycle. Thus, the signal current in this winding is ac. Although the input signal causes base current in both transistors, C_1 is charged by the base current of one transistor, and discharged by the base current of the other. Therefore, capacitive input coupling can be used with the circuit of Figure 10 even though it is operated in Class B.

Direct Coupled Circuit

Figure 11 shows a push-pull circuit that eliminates the need for an output transformer. The circuit can be constructed with either NPN or PNP transistors and does require a phase inverter to supply the input. A transformer is often used as the phase inverter as shown in Figure 11. This output circuit is often used in hi-fi equipment since it eliminates the need for an output transformer. Resistors R_1 and R_2 forward bias the base-emitter junction of Q_1 . Resistors R_3 and R_4 forward bias the base-emitter junction of Q_2 . The values for these bias resistors are chosen so that Q_1 and Q_2 conduct exactly the same when no input signal is applied to them. Thus, the V_{CC1} voltage is across Q_1 and the V_{CC2} voltage is across Q_2 . Points A and B are at the same potential and therefore no dc voltage is across the coil of the speaker. Thus, Q_1 , Q_2 , V_{CC1} and V_{CC2} form a bridge circuit. Since it is often difficult to maintain this balanced condition, a high value electrolytic coupling capacitor is placed in series with the speaker to eliminate any dc current through the speaker.

Input transformer T_1 supplies signals which are 180° out-of-phase to Q_1 and Q_2 . When the input signal aids the forward bias on Q_1 , current due to Q_1 conduction through the speaker coil increases. At the same time, Q_2 is driven toward cutoff and current due to Q_2 decreases in the coil. The circuit is operating push-pull, very much like the complementary-symmetry circuit of Figure 10.

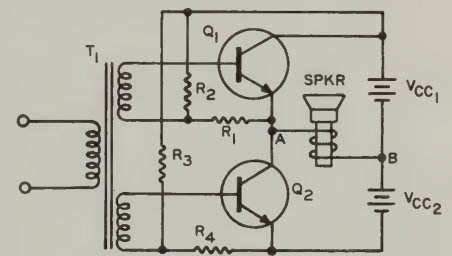


Figure 11

HIGH-FREQUENCY AMPLIFIERS

Before we continue, let's review what is meant by a high-frequency or a low-frequency signal. The chart of Figure 12 shows several frequency ranges.

TRANSISTOR PUSH-PULL AND HIGH FREQUENCY AMPLIFIERS

The audio-frequency range extends from about 20 to 20,000 hertz. A range of frequencies such as this is often referred to as a **BAND** of frequencies. The audio-frequency band is commonly called the low-frequency band. Any signal which has its frequency in this band is usually referred to as an audio-frequency signal or a low-frequency signal.

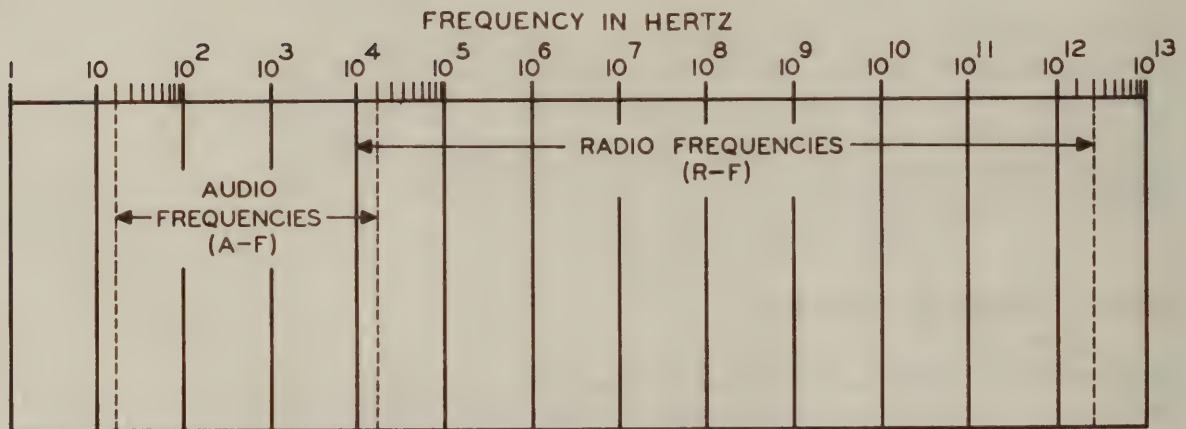


Figure 12

The radio-frequency band extends from about 10,000 to three-million million Hz (hertz). The radio-frequency band is commonly called the high-frequency band, although this band may be divided further into groups such as very high and ultra high frequency bands. Any signal with a frequency within this broad band generally is referred to as a radio-frequency signal. Notice that the audio and radio bands of frequencies overlap slightly. There is no sharp dividing line between the two bands. Frequencies which lie within this overlapped band may be referred to as either low-frequencies or high-frequencies, depending on the point of reference.

Tuned Amplifiers

Amplifiers can be classed as either **TUNED AMPLIFIERS** or **UNTUNED AMPLIFIERS**. A tuned amplifier contains at least one tuned or resonant circuit and usually amplifies only a certain band of frequencies. It rejects frequencies that are outside of this band. The tuned circuit or circuits are often used as part of interstage coupling networks between stages. Untuned amplifiers do not contain tuned or resonant circuits. However, even in untuned amplifiers the circuit part values may limit the useful range of frequencies amplified. Most low-frequency amplifiers are of the untuned type. Most high-frequency amplifiers are tuned amplifiers.

A resistance-capacitance coupled two-stage amplifier is shown in Figure 13. This is a common type of coupling used in many low-frequency amplifiers such as those found in hi-fi equipment. This circuit has certain factors which limit its use for very high and very low frequency signals.

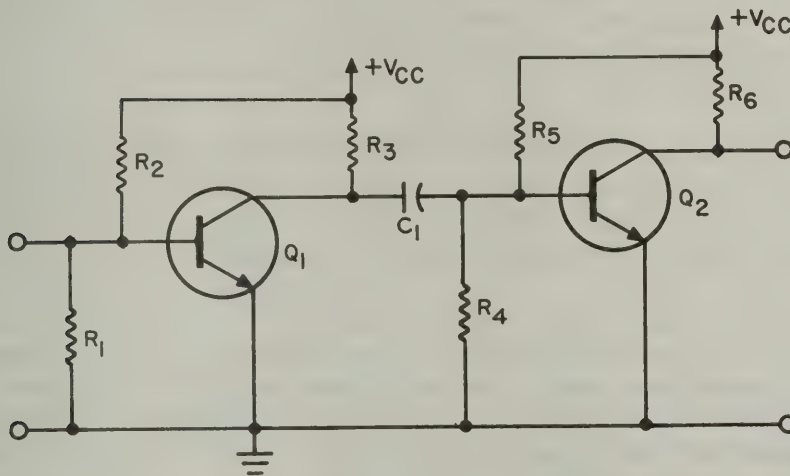


Figure 13

Briefly, the action of the coupling circuit between Q_1 and Q_2 is as follows. The Q_1 collector current variations cause corresponding voltage variations across R_3 . The voltage variations across R_3 cause corresponding variations of the Q_1 collector voltage with respect to ground. For voltage variations at the signal frequency, capacitor C_1 and resistor R_4 form a series circuit between the Q_1 collector and ground. The portion of this signal voltage across R_4 in this C_1R_4 voltage divider serves as the signal input for the Q_2 stage. Since the C_1 reactance varies inversely with frequency, the input to Q_2 developed across R_4 depends upon the signal frequency. With a given set of

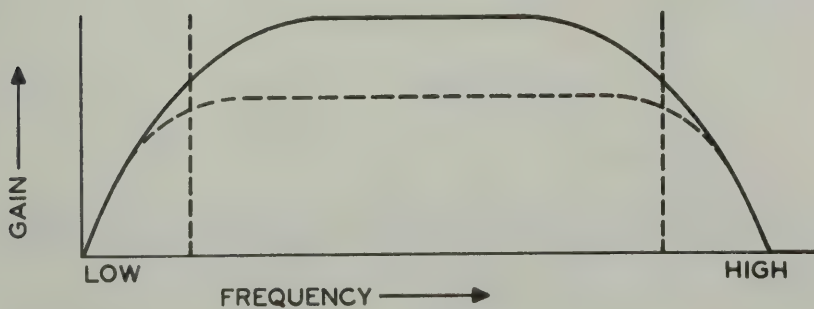


Figure 14

components, the C_1 reactance at some low-frequency is so high that all of the signal voltage is developed across the capacitor and none appears across the input of Q_2 . This is pictured by the gradual decrease in gain at the low-frequency end in the gain-versus-frequency graph of Figure 14. As the frequency increases, less voltage is developed across C_1 and more voltage appears at the input of Q_2 . Finally, at some medium frequency, the C_1 reactance is so low that practically all of the signal appears at the input of Q_2 and the gain of the two stages is high. This is shown by the middle area in the gain-versus-frequency graph of Figure 14.

In any transistor, INTERELECTRODE CAPACITANCE exists between each pair of electrodes, because electrically the electrodes are conductors separated by a dielectric. Also, in any circuit, stray capacitance exists between the various leads and from each lead to ground.

The collector-to-emitter interelectrode capacitance adds to the stray capacitance between the collector lead and ground to form a total known as OUTPUT CAPACITANCE (C_o). Also, the base-to-emitter interelectrode capacitance adds to the stray capacitance between the base lead and ground to form a total called INPUT CAPACITANCE.

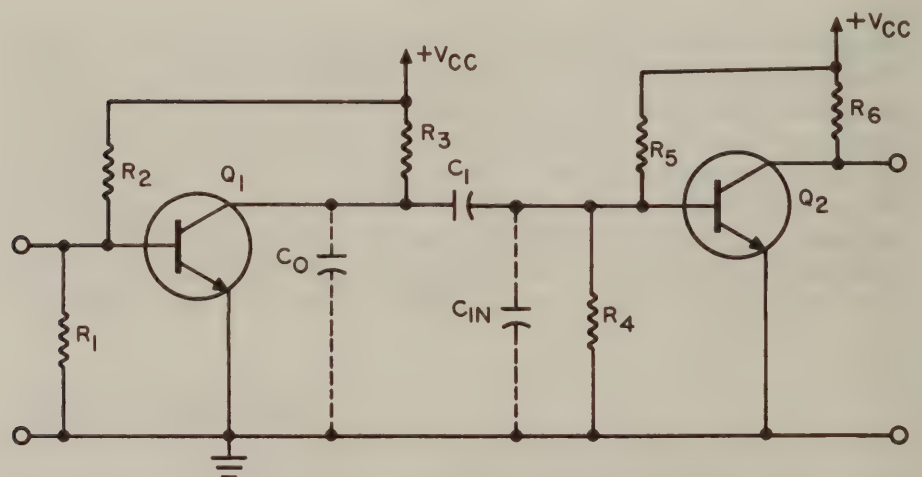


Figure 15

The electrical positions of input and output capacitance are illustrated in Figure 15, which is the same as Figure 13 except for the dashed line capacitor symbols C_o and C_{IN} . At the lower frequencies, C_o and C_{IN} have extremely high reactances. For this reason they can be ignored when we consider low-frequency operation as was done in Figure 13. At higher frequencies, the

reactances of C_O and C_{IN} decrease and they begin to shunt the signal to ground. C_O shunts the output signals of Q_1 to ground. C_{IN} shunts the input signals of Q_2 to ground. The result is a decrease in gain at high frequencies as shown by the solid line curve of Figure 14. At some high frequency, the C_O and C_{IN} reactances become so low that they short-circuit the signal and reduce the overall gain of the circuit to zero. For this reason, the circuit of Figure 13 and 15 is not a good high-frequency amplifier.

Shunt capacitances also affect the input to Q_1 and the output of Q_2 . However, at this point we are concerned only with the coupling between stages Q_1 and Q_2 . For this reason, we will disregard the shunting effects at the input and output of the circuit.

One method of improving the high-frequency response in the circuit of Figure 15 is to lower the collector load (R_3) resistance. Of course, this change lowers the output at all frequencies because it lowers the collector load impedance. The effect on frequency response is shown by the dashed line curve of Figure 14. It does, however, raise the frequency at which the shunt capacitance becomes a large part of the load impedance. Notice that the bend in the upper end of the dashed line of Figure 14 occurs at a higher frequency than does the bend of the solid line representing a higher collector load resistance.

We have shown that the effect of these shunt capacitances makes the resistance-capacitance coupled amplifier a poor high-frequency amplifier. However, in other types of circuits, these shunt capacitances can be used to actually improve the high-frequency amplification. Such a circuit is shown in Figure 16. Transistors Q_1 and Q_2 are impedance coupled. Inductor L_1 replaces resistor R_3 of Figures 13 and 15. Capacitors C_2 and C_3 have been added to bypass the collector supply. Without these capacitors, the high-frequency signal may cause a voltage to be developed across the internal impedance of the collector supply resulting in an undesired variation of collector voltages. C_2 and C_3 were omitted in Figures 13 and 15 for simplicity.

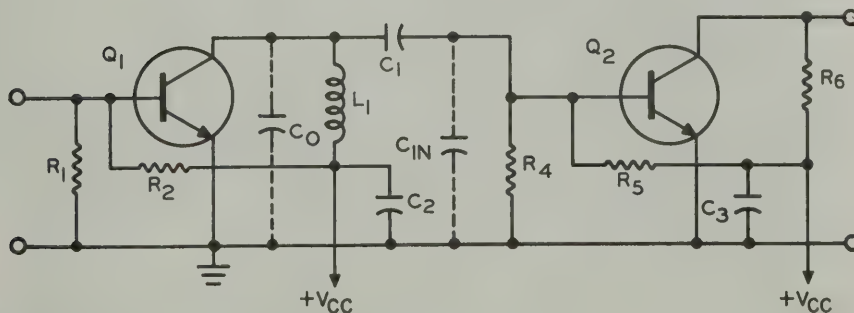


Figure 16

At high signal frequencies, the reactances of C_2 and C_3 are very low. C_1 may also be considered a short circuit for high-frequency signals. In effect, the output capacitance C_o of Q_1 and the input capacitance C_{IN} of Q_2 are in parallel with inductor L_1 . C_o , C_{IN} and L_1 form a tank circuit that resonates at some high frequency. The impedance of a parallel LC circuit rises as the frequency of the signal approaches the resonant frequency of the coupling circuit. This type of circuit has a high impedance at or near the resonant frequency. Thus, if the signal frequency is near the resonant frequency of the coupling circuit, Q_1 collector current through the tank develops a large output which is applied to the input of Q_2 .

To give us greater control of the resonant frequency of the tank circuit, we generally place a capacitor across L_1 . Such a capacitor is shown as C_4 in Figure 17. The tank circuit capacitance is now the sum of parallel capacitances C_4 , C_o and C_{IN} . This total capacitance and L_1 form the tank circuit. The value of C_4 is chosen to make the tank circuit resonate at the desired frequency.

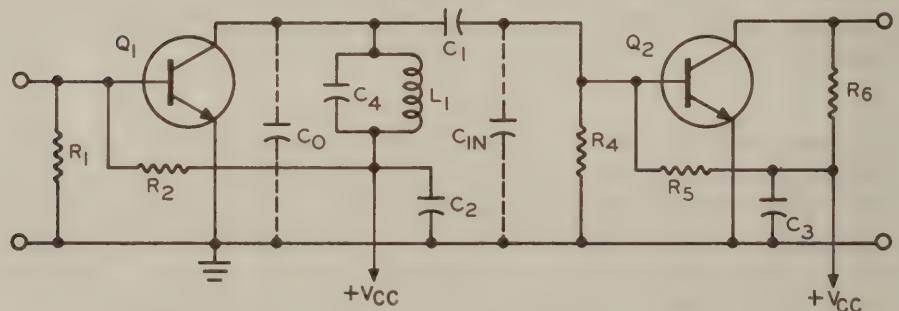


Figure 17

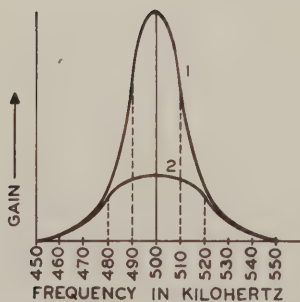


Figure 18

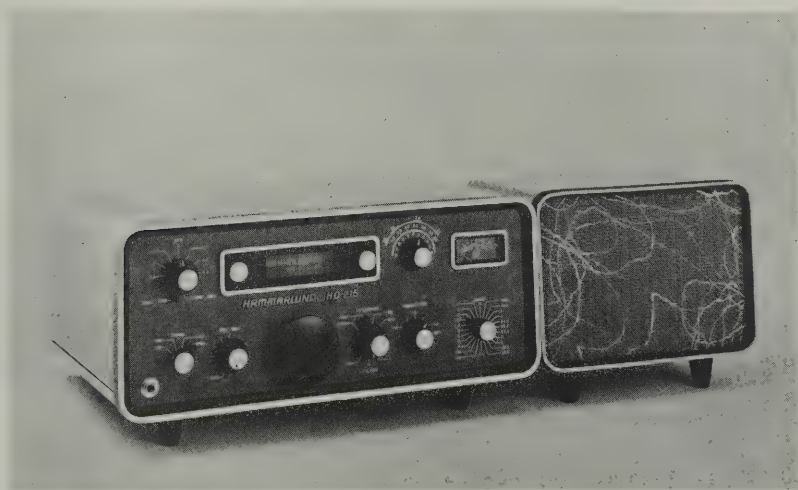
Since the coupling circuit of Figure 17 contains one tuned circuit, the amplifier is a **SINGLE-TUNED AMPLIFIER**. In a parallel tuned circuit, the inductive reactance predominates at frequencies below resonance, and the capacitive reactance predominates at frequencies above resonance. These reactances vary with frequency, so that the total impedance of the parallel circuit varies with signal frequency. The impedance, Z_L , of the tank is maximum at the resonant frequency of the tuned circuit. Z_L falls off for frequencies above and below resonance. Gain of the two stages then varies with frequency and is maximum at resonance because Q_1 develops its output across a higher impedance at the resonant frequency.

The curves of Figure 18 show how the amplification varies, depending upon the Q of the tuned circuit. In both cases the amplification is maximum at

a frequency of 500 kHz, and lower at frequencies above and below this point. In any tuned amplifier, collector load impedance Z_L depends upon the Q or figure of merit of the tuned circuit as well as upon the frequency. The Q depends on the relationship between the reactance and resistance of a circuit. Resistance shunting a parallel tuned circuit reduces the Q or figure of merit. In Figure 18, curve 1 represents the response of an amplifier in which the tuned circuit has a relatively high Q , while curve 2 represents the response of a low- Q circuit. Due to the shunting effects of bias resistor R_4 , the internal collector resistance of transistor Q_1 and the internal base resistance of Q_2 in Figure 17, the effective Q of the tuned circuit always is less than the Q of the coil. Remember that the effective Q of a coil is the ratio of the coil's reactance to its SERIES resistance, $Q = X_L/R_{\text{SERIES}}$. On the other hand, the effective Q of a tuned circuit depends upon the shunt resistance and is the ratio of the shunt resistance to the coil reactance, $Q_{\text{eff}} = R_{\text{SHUNT}}/X_L$.

Most signals include a range or band of frequencies. The curves of Figure 18 indicate maximum amplification at the resonant frequency. In practice, it is considered that all output signals with magnitudes equal to or greater than 70.7% (-3db) of the resonant frequency output are amplified properly. This includes signals having frequencies both above and below the resonant frequency. All these frequencies are said to be "passed" by the amplifier. They cover a range that is known as the passband. Depending upon the width of the passband, in a given application an amplifier may be considered to have narrow, medium, or wide frequency response or bandpass.

The bandpass of a tuned amplifier depends upon the Q of the tuned circuit. The dashed lines in Figure 18 indicates a bandpass from 490 to 510 kHz



The bandpass of the tuned amplifier circuits in a communications receiver is very important. Excessive bandwidth will introduce adjacent channel interference. If the bandwidth is too narrow, the signals will not be received properly.

Courtesy Hammarlund Manufacturing Co.

for a high-Q circuit with the response of curve 1, and a bandpass from 480 to 520 kHz for a lower-Q circuit with the response of curve 2. Thus, the lower Q permits a broader bandpass but results in less amplification.

The response curve of the circuit of Figure 17 is sharply pointed when a relatively high-Q resonant circuit is employed. Although this type of response is satisfactory for many applications, there are others such as in broadcast and FM radio receivers which require a steep-sided, relatively flat-topped response characteristic. The flat-topped response provides substantially constant amplification for all the desired signal frequencies, and at the same time sharply rejects frequencies outside the desired band.

A steep-sided, flat-topped response may be obtained by using a form of transformer coupling as shown in Figure 19.

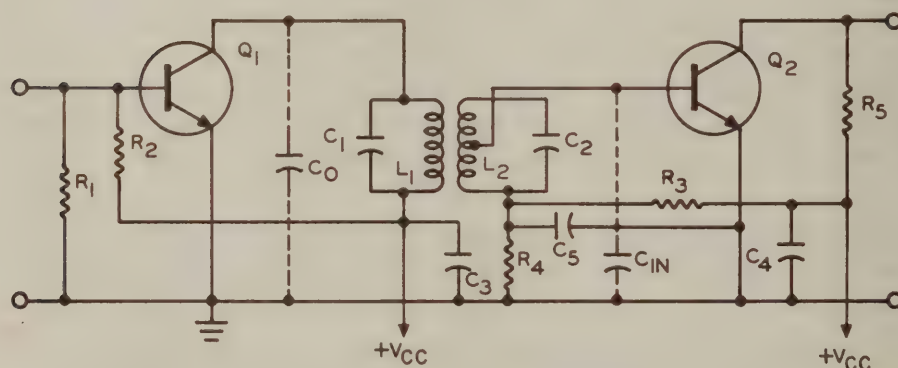


Figure 19

The inductors which are used as primary and secondary of the coupling transformer are each part of a tank circuit. One tank circuit is the collector load for the first stage. The other tank circuit is in the base circuit of the second stage. Since there are two tuned circuits, the amplifier is called a **DOUBLE-TUNED AMPLIFIER**. Capacitor C_1 in parallel with the output capacitance of Q_1 tunes transformer primary L_1 , while secondary L_2 is tuned by the parallel combination of C_2 and C_{IN} .

Notice that the tap on L_2 is connected to the base of Q_2 . The tapped arrangement is used to provide an impedance match between the coil and the Q_2 input. The Q_2 base circuit is a low impedance circuit whereas the tuned transformer secondary has a high impedance. The Q_1 collector can also be connected to a tap on the primary to provide additional matching. Imped-

ance matching is explained later in this lesson. The purpose of C_5 in Figure 19 is to pass the signal frequencies from the bottom of the tank circuit directly to the emitter of Q_2 . This prevents a signal loss across bias resistor R_4 .

With the two tuned circuits, primary and secondary, resonant to the same frequency, the frequency response depends partly on the degree of coupling between the coils. When the degree of coupling is small, the interaction between the primary and secondary is low and the amplification curve resembles that of ordinary resonance, as shown by curve 1 in Figure 20.

As the coupling is increased to the CRITICAL POINT, the response will become peaked as shown by curve 1 in Figure 18. When the coupling is increased beyond the critical point, the curve takes on a double-humped shape as shown by curve 2 in Figure 20. A further increase of coupling does not increase the amplification, but causes the two humps to spread apart and the valley between them to deepen. Again, the shape of the amplification curve depends on the Q's of the tuned circuits, being higher and steeper sided for higher-Q circuits.

Thus, an amplifier employing the circuit of Figure 19 may be adjusted to provide essentially uniform output for signals having frequencies within a desired band, but little output for signals having frequencies outside the band limits.

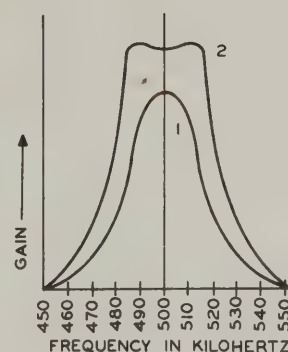


Figure
20

Neutralization

Figure 21 shows a transistor high-frequency amplifier circuit using a PNP transistor. The circuit is very similar to those which you have already studied. R_1 and R_2 apply forward bias to the base-emitter junction of Q_1 . In this amplifier, the collector load is the tank circuit composed of C_2 and the primary of T_2 . The output signal is inductively coupled from the tank circuit to the secondary of T_2 . The input signal is applied to transistor Q_1 by transformer T_1 . The tap of T_1 is connected directly to the base of Q_1 . The bottom of T_1 is connected to the emitter of Q_1 through capacitor C_1 . The reactance of C_1 is very small at signal frequencies. Thus, C_1 passes the input signal easily but does not affect the dc bias voltages.

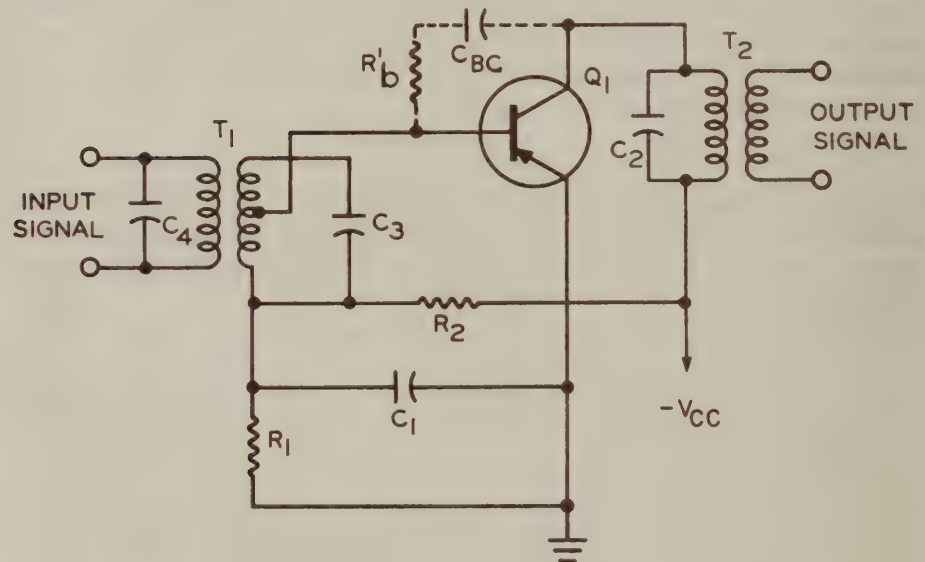


Figure 21

In this figure, dashed line capacitor symbol C_{BC} represents the internal capacitance of the base-collector junction of the transistor. R'_b represents the resistance of the base-collector junction. Keep in mind that these are the internal resistance and capacitance of the transistor. They are always present within a transistor. Notice carefully the represented locations of C_{BC} and R'_b . They are connected from the output of the transistor, the collector, to the input of the transistor, the base. At low-frequencies, the reactance of C_{BC} is quite large. C_{BC} thus offers a large opposition to any signal which attempts to pass through it in either direction. Thus, at low signal frequencies, the effect of C_{BC} can usually be ignored.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

PUSH-PULL TRANSISTOR AMPLIFIERS

LIMITATIONS OF RC COUPLING

COMPLEMENTARY-SYMMETRY CIRCUIT

DIRECT COUPLED CIRCUIT

HIGH-FREQUENCY AMPLIFIERS

TUNED AMPLIFIERS

14. With the circuit of Figure 7 operating Class B, the voltage drops developed across C_1 and C_2 provide forward bias. True or False?
15. How does a complementary-symmetry circuit differ from a conventional push-pull amplifier?
16. In the circuit of Figure 10, as the input signal goes positive, Q_1 is driven into (cutoff) (conduction) and Q_2 is driven into (cutoff) (conduction).
17. For proper operation, R_1 in the circuit of Figure 10 should have a high value. True or False?
18. Does the circuit of Figure 11 require 180° out of phase input signals?
19. With no signals applied to the circuit of Figure 11, does the speaker carry any dc current?
20. How does a tuned amplifier differ from an untuned amplifier?
21. Explain how the value of the coupling capacitor in an RC coupled amplifier controls the low-frequency response.
22. Does the shunt capacity in a circuit affect the high or low-frequency response?
23. Reducing the value of the collector load improves the high-frequency response. True or False?
24. What is the name given to the circuit of Figure 17?
25. Why does the circuit of Figure 17 provide high gain at one frequency?
26. In Figure 18, curve 1 is for an amplifier stage with a low effective Q. True or False?

27. Increasing the circuit Q (increases) (decreases) bandwidth.
28. What happens to the circuit Q as the value of shunting resistance is decreased?
29. Why is the base of Q_2 connected to a tap on L_2 in Figure 19?
30. If you desired a double tuned amplifier with a wide frequency response, would you employ over or under coupling?
31. What would happen to the response of the circuit of Figure 19 if the base of Q_2 were connected to the top of L_2 rather than to the tap?

At higher frequencies, signals developed at the collector of Q_1 find a path through C_{BC} back to the base. This is called **FEEDBACK**. At high-frequencies the reactance of C_{BC} is low and offers little opposition to any signal which attempts to pass through it. Large amounts of feedback may be applied from the collector of Q_1 through R'_b and C_{BC} to the base of Q_1 . This backward transfer of energy is the feedback that can cause difficulties at high-frequencies.

The feedback may add to the input signal or it may oppose the input signal, depending on the phase shifts in the circuit. If the feedback aids the input signal the result is called **POSITIVE FEEDBACK** or **REGENERATIVE FEEDBACK**. This causes the circuit to be unstable. If the feedback opposes the input signal the result is called **NEGATIVE FEEDBACK** or **DEGENERATIVE FEEDBACK**. This condition decreases the gain of an amplifier.

One way to prevent unwanted feedback is through **NEUTRALIZATION**. A circuit is neutralized by introducing an additional voltage, equal to the feedback voltage but exactly 180° out of phase with it. This cancels or neutralizes the feedback voltage. Neutralization is often called **UNILATERALIZATION** in a transistor circuit since the effects of both resistance and capacitance must be eliminated.

Figure 22 shows the amplifier of Figure 21 with a neutralizing circuit added. A resistor, R_N , is added in series with neutralizing capacitor C_N . The primary winding of T_2 has a center tap which connects to the negative terminal of

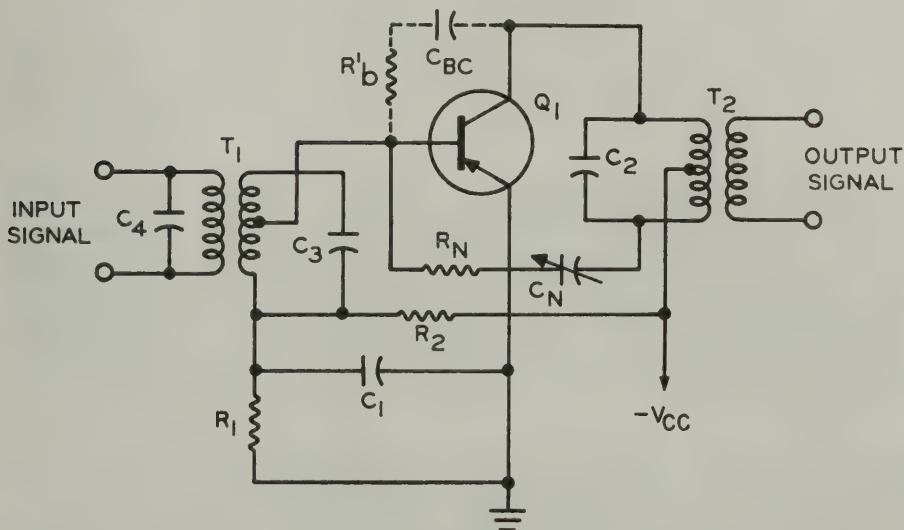


Figure 22

$-V_{CC}$. Because $-V_{CC}$ has low impedance at signal frequencies, the center tap of the collector tank is at signal ground potential. The signal voltage at the lower end of the tank is 180° out of phase with the signal voltage at the upper end, which is connected to the collector.

The collector signal voltage produces one feedback voltage at the base due to C_{BC} and R'_b . The voltage across the lower half of the collector tank produces a second feedback voltage at the base due to C_N and R_N . C_N is adjusted until the amplitude of the neutralizing voltage is equal to that of the undesired feedback voltage. These two voltages are 180° out of phase, and therefore cancel each other. In some cases R'_b is negligible and only C_N is used. In either case, when the unwanted feedback is completely canceled, the circuit is said to be neutralized.

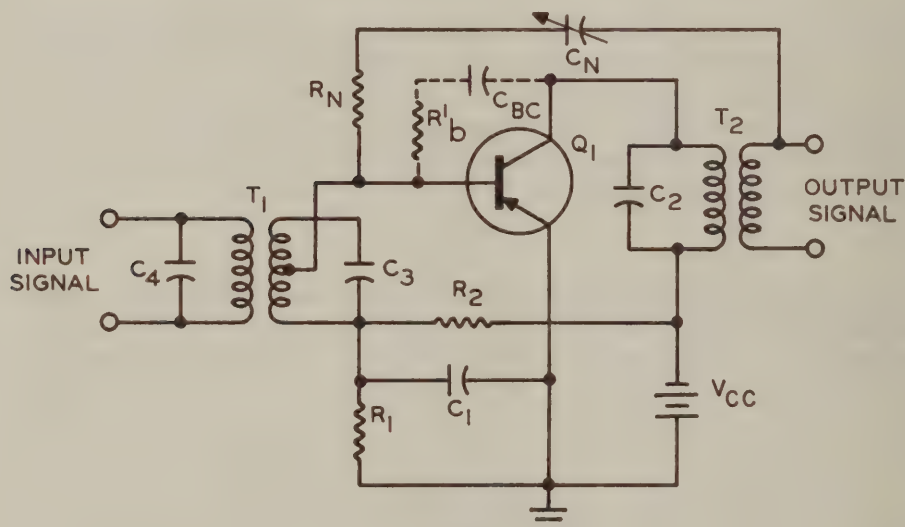


Figure 23

Figure 23 shows a second method of neutralizing a transistor high-frequency amplifier. In this case, the collector tank coil does not have a center tap. The needed out-of-phase voltage is obtained from the secondary winding of T_2 . C_N is connected to the end of this winding at which the signal voltage is 180° out-of-phase with the Q_1 collector signal. Hence, the current fed back through R_N and C_N is out-of-phase with that fed back through R'_b and C_{BC} .

Push-pull and Parallel H-F Amplifiers

To obtain greater power output, push-pull or parallel operation can be used in high frequency amplifiers. Push-pull operation has the advantage of can-

celing even harmonic distortion. Parallel operation has the advantage of not requiring a phase inverter arrangement.

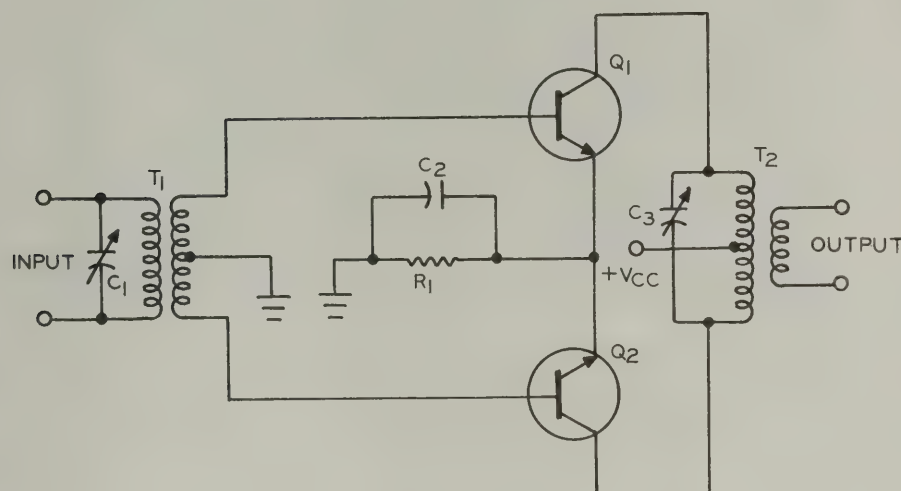


Figure 24

Figure 24 shows a push-pull high-frequency amplifier. The required 180° out-of-phase input signals are provided by employing a center tapped secondary on T_1 . Note that the secondary of T_1 is untuned. Since the base circuit impedances are low, the resulting Q is so low that tuning is impractical. In this circuit Q_1 and Q_2 are operated with zero bias, and thus operate Class B. Power amplifiers are generally operated Class B or C to obtain high efficiency. Voltage amplifiers are operated Class A to obtain low distortion. Remember that the resonant load circuit develops a sinusoidal output from the collector current pulses. If Class C operation is desired, a fixed bias supply could be employed.

Thermal runaway protection for Q_1 and Q_2 is provided by R_1 and C_2 in the emitter circuit. The reverse bias provided by R_1 provides the stability. This arrangement is generally employed in high-frequency power amplifiers since the dc resistance of the load circuit is low. Depending upon the frequencies involved, a neutralization circuit might be employed in the circuit of Figure 24. Either arrangement: Figure 22 or 23 could be employed.

A parallel amplifier arrangement is shown in Figure 25. The base and collector circuits are connected in parallel. Again, note that the secondary of T_1 is untuned. Due to the parallel connection, the collector impedance is low so the collector is connected to a tap on the primary of T_2 . Class C

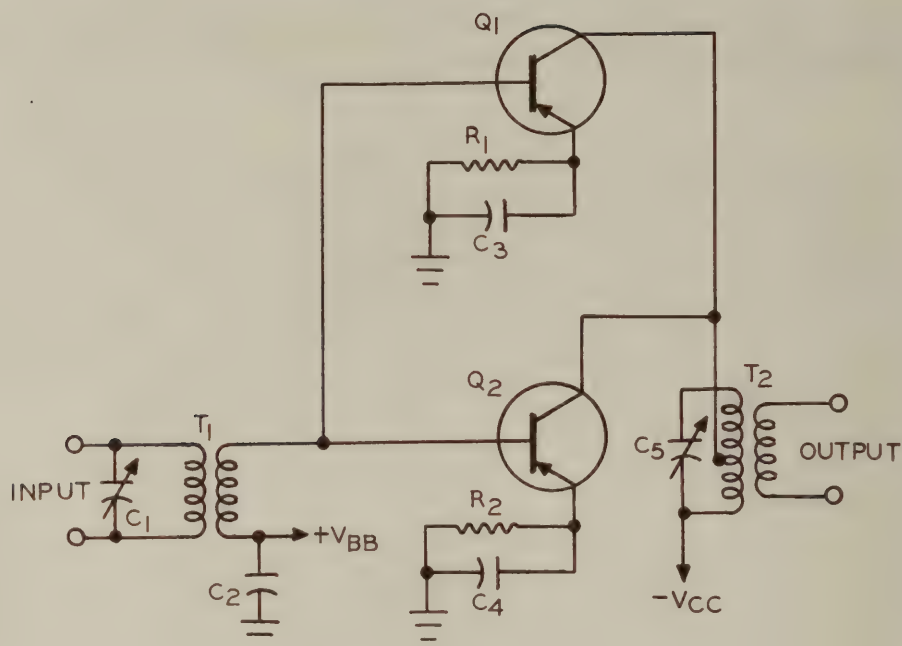


Figure 25

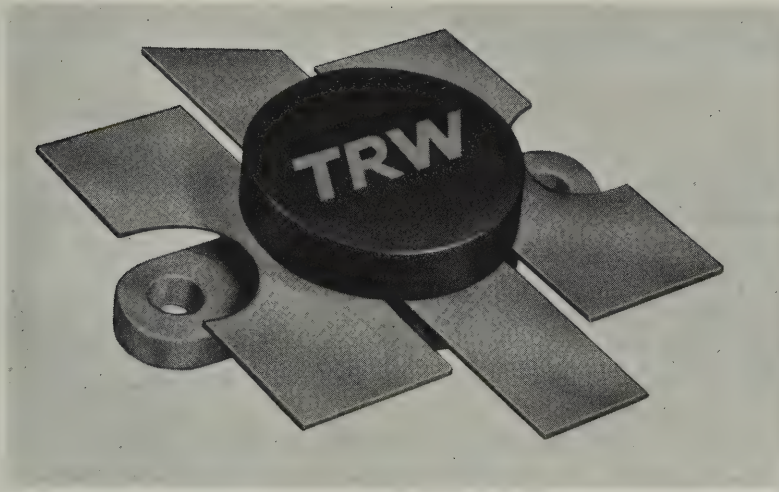
operation is provided by the fixed bias supply $+V_{BB}$. Thermal runaway protection is provided by R_1C_3 and R_2C_4 .

As in the push-pull circuit, the neutralization circuits of Figure 22 or 23 can be added to the circuit of Figure 25. For best operation in the circuits of Figures 24 and 25, the transistors should be matched. In some circuits, the bias on each transistor is adjusted separately to compensate for differences.

As in vacuum tube circuits, a transistor h-f amplifier can be operated as a frequency multiplier. The collector circuit is tuned to the desired harmonic of the input signal frequency. Class B or C operation is employed to obtain the desired harmonic content in the collector circuit. Since the input and output circuits do not operate at the same frequency in a frequency multiplier, neutralization is not required.

High-frequency Components

At high frequencies, ordinary iron and steel cannot be used for coil cores. This is mainly because losses from currents induced in the core are so great at high-frequencies. These induced currents and the resulting losses increase with frequency.



This power transistor is especially designed for r-f power amplifier service in the 225 - 400 MHz range. Although physically the unit appears as a single transistor, it is electrically several transistors in parallel.

Courtesy Semiconductor Division of TRW Inc.

Coils for high-frequency circuits either have no magnetic cores or have special magnetic cores which reduce losses. The coils without magnetic cores are called air-core coils. Magnetic cores used in high-frequency coils are made from finely divided or "powdered" iron. The iron particles are held together by an insulating binding material. Each particle is insulated from the others. This prevents the flow of current between particles, and thus reduces losses.

Current tends to flow in the outer part of the wire that carries it. This is called SKIN EFFECT. In the case of dc and low-frequency ac, the skin effect is slight. But with high-frequency ac, almost all the current is carried in the "skin" of the wire. This reduces the useful cross-section of the wire. It is the same as reducing the size of the wire because it increases the wire's resistance. At high frequencies a given wire has several times as much resistance as it has for dc. To avoid confusing it with the dc resistance, this high resistance is referred to as the r-f (radio-frequency) resistance. Various types and sizes of wire are used for r-f coils, depending on the frequency of the signal.

IMPEDANCE MATCHING

Any device which generates or amplifies a signal works best when it delivers its output to a load of a certain impedance. This impedance may be a pure resistance or it may be a combination of resistance and reactance. For example, a transistor usually requires a fairly high impedance in its collector

circuit if it is to develop the desired output signal. The transistor stage may not function properly if this impedance is not provided.

On the other hand, a dynamic speaker is an example of a device that places a very low impedance across the circuit to which it is connected. If we connect a dynamic speaker to the collector circuit of a transistor, a **MISMATCH** exists. That is, the speaker impedance prevents the transistor from operating properly if connected in this way. There are devices, however, which can be placed between the speaker and transistor to make the speaker impedance **APPEAR** to be that which the transistor collector circuit requires. This would then provide impedance matching.

As another example, consider the transformer in the circuit of Figure 19. The tuned secondary of this transformer presents a high impedance between its terminals at the resonant frequency of the secondary circuit. A low impedance placed across the secondary would change the bandpass of the amplifier. The input or base circuit of transistor Q_2 , however, places a low impedance across the circuit it is connected to. The impedance from a tap on the transformer to one end is less than the impedance of the full winding. Therefore, by tapping the transformer secondary, a point may be found where the impedance is less than that of the full tuned secondary and about equal to that of the base circuit of Q_2 . Thus, the base resistance of transistor Q_2 does not appear across the tuned transformer in such a way that will prevent it from functioning properly.

Impedance matching may be necessary to provide power transfer, prevent distortion, or to maintain proper bandpass, among other reasons. A number of devices can be used to obtain a match of impedances. Some of these are shown in the following part of this lesson.

OUTPUT TRANSFORMER COUPLING

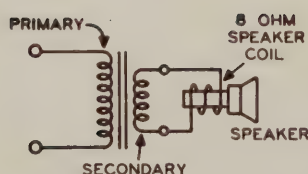


Figure
26

To study the use of the output coupling transformer, let's take a practical example. Suppose we have a speaker which has a voice coil impedance of 8 ohms. We also have an amplifier designed to work into an output impedance of 800 ohms. To obtain maximum power output, we must match these impedances. This can be done by using the transformer T_1 as shown in Figure 26. The primary of the transformer is connected to the amplifier, and the secondary is connected to the 8 ohm speaker coil.

Suppose the transformer of Figure 26 has a 10 to 1 turns ratio. This means there are 10 times more turns on the primary than on the secondary. The voltage across the secondary is then 1/10 that across the primary. How-

The following Practice Exercise questions cover the subjects which you have just studied. They are:

HIGH-FREQUENCY AMPLIFIERS

NEUTRALIZATION

PUSH-PULL AND PARALLEL H-F AMPLIFIERS

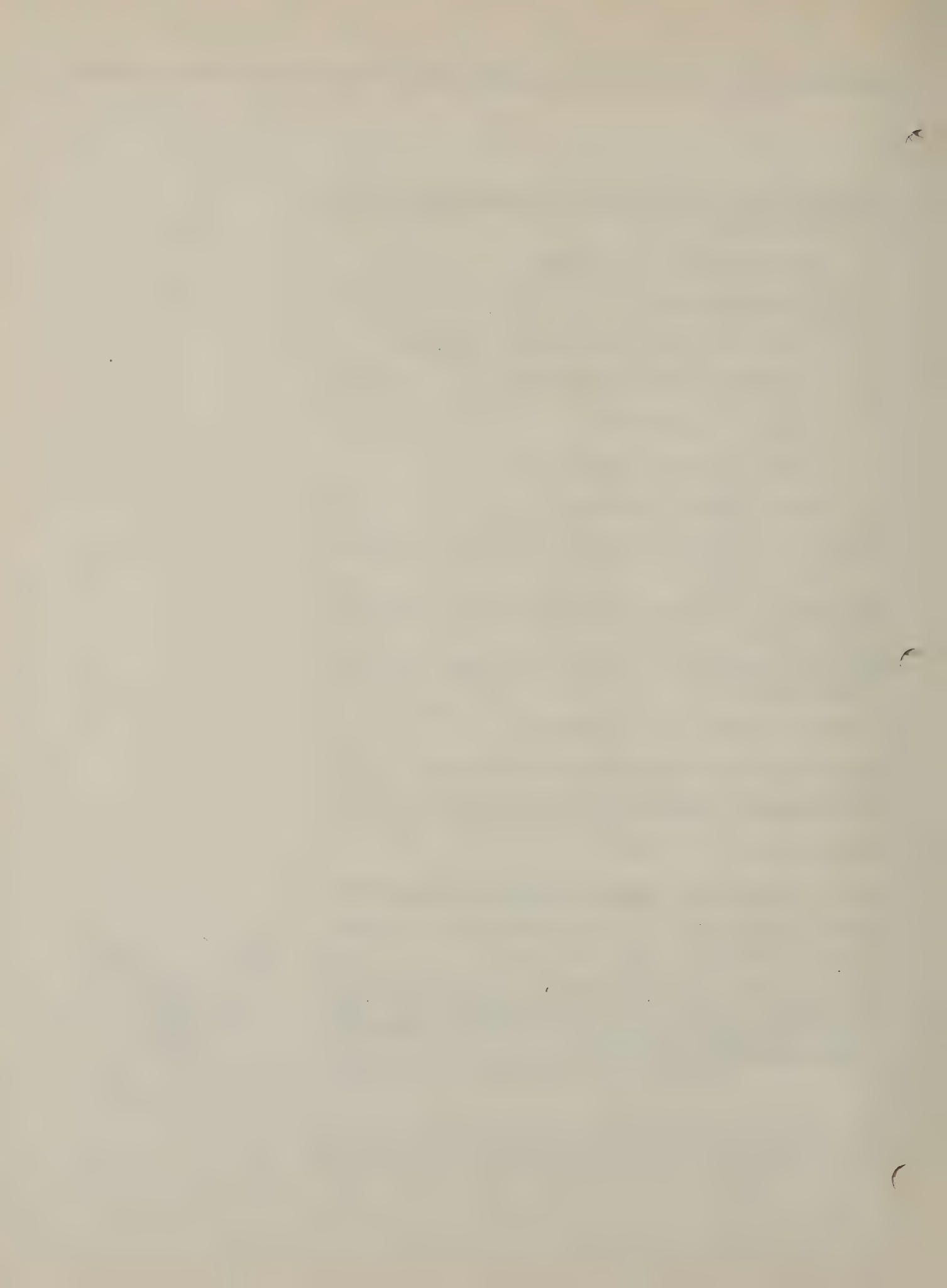
HIGH-FREQUENCY COMPONENTS

IMPEDANCE MATCHING

OUTPUT TRANSFORMER COUPLING

TUNED OUTPUT COUPLING

32. When do the effects of the interelectrode capacitance of a transistor become important?
33. In Figure 22, the signal at the bottom of the T_2 primary is in phase with that at the top. True or False?
34. What are some of the advantages of push-pull operation over single-ended operation?
35. Zero bias provides what class of operation?
36. Why would a power amplifier be operated Class B or C?
37. What material is generally used for cores in high-frequency coils?
38. What is meant by "skin effect"?
39. Why must an impedance match between source and load be obtained?
40. How does the impedance at the primary and secondary of an untuned transformer compare to the turns ratio?
41. Suppose that in the arrangement of Figure 26, the T_2 secondary is connected to a speaker that has a 4 ohm impedance, and the circuit connected to the primary must work into an impedance of 1000 ohms. What is the required turns ratio?



ever if the speaker uses all the power delivered to the transformer, the current drawn from the secondary must be 10 times greater than that in the primary. Thus, the amplifier provides output power in the form of high voltage and low current. This power is delivered to the speaker in the form of low voltage and high current. Ohm's Law, $R = E/I$, relates voltage and current. If the voltage E is high and the current I is low, the resistance (or impedance) appears high. This is the situation at the transformer primary. The impedance presented to the amplifier appears much higher than the 8 ohms of the speaker. The impedance presented by the primary of an untuned transformer is the **SQUARE OF THE TURNS RATIO TIMES THE IMPEDANCE CONNECTED TO THE SECONDARY**. In this example the turns ratio is 10. The square of the turns ratio is 100. Multiplying by the speaker impedance of 8 ohms gives a primary impedance of 800 ohms. Thus, the amplifier and speaker are matched.

TUNED OUTPUT COUPLING

As you know, the stages in high-frequency amplifiers are usually coupled by tuned circuits. This is done to obtain a large signal gain with the desired bandwidth. The same is also true in the output stage when a load is coupled to the output of the amplifier.

As an example, let's examine the output circuit that is commonly used with a high-frequency radio transmitter. In Figure 27, the primary of transformer T_1 and capacitor C_1 form a tank circuit which is connected in the final stage of the high-frequency amplifier. The parallel tuned circuit of the primary is basically a high impedance circuit. Power delivered to this circuit is in the form of a high voltage, low current signal. However, a dipole antenna places a low impedance across a source when operated at its resonant frequency.

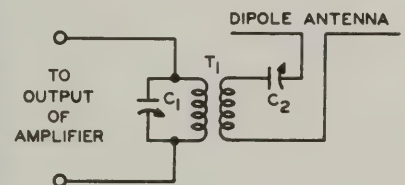


Figure 27

In Figure 27, we see that the impedance of the dipole is connected in series with C_2 and the transformer secondary. We may consider the secondary inductance, capacitor C_2 and the antenna impedance, as a series circuit. In this series resonant circuit, the current is high but the voltage across the antenna impedance is low. Thus, the power to the antenna is in the form of low voltage and high current. Here we have matched the low impedance of the antenna to the high impedance of the amplifier. In matching with tuned circuits, the calculations are more involved than with audio transformers. The Q of each circuit, the amount of coupling, the inductance of each winding and the bandwidth must all be considered.

SUMMARY

Push-pull operation is employed if increased power output is desired. This type of operation also eliminates some forms of distortion. Most transistor push-pull amplifiers require a device which can apply 180° out-of-phase signals to the inputs of the two transistors. A phase inverter or a tapped input transformer may be used to provide these out-of-phase signals. The complementary-symmetry circuit is a push-pull transistor amplifier which does not require two signals at its input. This circuit employs both an NPN and a PNP transistor.

Many high-frequency amplifiers are tuned amplifiers. Tuning is provided by resonant circuits which make up interstage coupling networks. The input and output shunt capacitances of the amplifier are used as part of the resonant circuits. When the coupling network contains one resonant circuit, the method is called single-tuned coupling. Double-tuned coupling uses two resonant circuits. At high frequencies, the collector-to-base interelectrode capacitances of the transistor can no longer be ignored. The reactances of these capacitances at high frequencies are low enough that energy from the output of the amplifier is fed back to the input of the amplifier through the internal capacitances of the transistor. Neutralization is used to prevent undesired effects of this feedback.

To obtain the best amplifier operation it is necessary for the amplifier load impedance to be of the proper value. Where the load is not the value desired, it is possible to make the load impedance appear different by applying the output signal to some device such as a transformer. It is thus possible to change the voltage-current relationship to match the requirements of the amplifier. This is known as impedance matching.

Various types of coupling methods are used to couple the output of an amplifier to a load device. The type of coupling used depends on the type of amplifier and the type of load. Common types of output coupling are transformer coupling, tuned transformer coupling and direct coupling.

IMPORTANT DEFINITIONS

COMPLEMENTARY-SYMMETRY CIRCUIT — A push-pull transistor amplifier employing a PNP and an NPN transistor that does not require 180° out-of-phase input signals.

CROSSOVER DISTORTION — Waveform distortion produced in a push-pull amplifier due to operation on the low current portion of the characteristic curves.

DOUBLE-TUNED AMPLIFIER — An amplifier in which the coupling network has two tuned circuits.

NEUTRALIZATION — A method of canceling the effects of feedback by introducing a second feedback path around an amplifier. Generally called unilateralization in a transistor amplifier.

PHASE INVERTER — A stage that supplies two signals 180° out-of-phase with each other.

SINGLE-TUNED AMPLIFIER — An amplifier in which the coupling network has only one tuned circuit.

SPLIT-LOAD PHASE INVERTER — A phase inverter where one output is taken from the collector load resistance and the other output from the emitter load resistance.

TUNED AMPLIFIER — An amplifier which uses tuned circuits as coupling networks.

UNTUNED AMPLIFIER — An amplifier which does not contain tuned or resonant circuits in its coupling networks.

PRACTICE EXERCISE SOLUTIONS

1. Generally, a push-pull amplifier requires two input signals 180° out-of-phase.
2. Crossover distortion occurs in a Class B push-pull amplifier due to the fact that bends occur in the $I_C - I_B$ curves at the zero bias points.
3. Crossover distortion is eliminated by applying a small amount of forward bias to each transistor. This aligns the $I_C - I_B$ transfer curves as shown in Figure 4.
4. False — Push-pull operation cancels even harmonic distortion, not odd harmonic distortion.
5. reverse — The emitter resistors, R_1 and R_2 , provide a small amount of reverse bias for Q_1 and Q_2 . This reverse bias is much smaller than the forward bias due to R_3 and R_4 . The net result is that the bases are forward biased.
6. The net current in the output transformer primary is zero under no signal conditions.
7. The output power of a push-pull amplifier is twice that of a single-ended stage.
8. The low-frequency response is limited by the inductance of the primary winding. Distributed capacity and leakage inductance limit the high-frequency response.
9. The circuit of Figure 8 is referred to as a split load phase inverter.
10. less, more — When the input signal drives the base of Q_1 more positive, collector current decreases. This reduces the voltage drop across the emitter resistor and increases the Q_1 collector-to-emitter voltage.
11. No, with a large value of emitter resistance that is unbypassed, the gain is low.
12. The signal at the base of Q_2 is 180° out-of-phase with the input signal. This is due to the phase shift introduced by the Q_1 stage.
13. Yes, both amplifier circuits in Figure 9 are common-emitter amplifiers.
14. False — The voltage developed across the coupling capacitors provide reverse bias, shifting the Class of operation from B to C.
15. A complementary-symmetry amplifier employs a PNP and an NPN transistor. As a result, it does not require a phase inverter circuit.

PRACTICE EXERCISE SOLUTIONS (Continued)

16. cutoff, conduction — Since Q_1 is a PNP transistor, the positive input signal drives it into cutoff and since Q_2 is an NPN transistor, the positive input signal drives it into conduction.
17. False — R_1 must have a low value so that the same amplitude input signal is applied to both Q_1 and Q_2 .
18. Yes, transformer T_1 serves as a phase inverter supplying the necessary 180° out-of-phase input signals.
19. No, if the circuit is properly adjusted, the voltage between points A and B is zero.
20. A tuned amplifier provides amplification for a narrow range of signal frequencies while an untuned amplifier is not frequency selective.
21. As the signal frequency decreases, the reactance of the coupling capacitor increases. This causes most of the signal voltage to appear across the capacitor, thus reducing the signal delivered to the next stage. The value of the coupling capacitor is chosen so that its reactance is low at the lowest signal frequency to be amplified.
22. The shunt capacity affects the high-frequency response.
23. True — Reducing the load resistance reduces the shunting affect of the shunt circuit capacity.
24. The circuit of Figure 17 is called a single-tuned amplifier.
25. In the circuit of Figure 17, the load impedance is a parallel LC or tank circuit. This circuit has a high impedance at its resonant frequency, thus providing maximum gain at this frequency.
26. False — Curve 1 is for an amplifier with a high effective Q and curve 2 is for an amplifier with a low effective Q.
27. decreases — Increasing the Q decreases bandwidth and decreasing the Q increases bandwidth.
28. Decreasing the value of the shunting resistance decreases Q and increases bandwidth.
29. The base of Q_2 is connected to a tap on L_2 to provide an impedance match since the tank circuit has a high impedance and the base circuit has a very low impedance.
30. Over coupling would be employed to provide a wide frequency response.

PRACTICE EXERCISE SOLUTIONS (Continued)

31. Connecting the base of Q_2 to the top of L_2 would shunt the tank circuit with a low impedance. This would reduce circuit Q , reducing gain and increasing bandwidth.
32. The interelectrode capacity of a transistor becomes important at high frequencies where its reactance decreases.
33. False — Due to the center tapped arrangement of T_2 , the signal at the bottom is 180° out-of-phase with the signal at the top.
34. Push-pull operation provides increased power output and cancellation of even harmonic distortion.
35. Zero bias provides Class B operation.
36. A power amplifier is operated Class B or C to obtain high efficiency.
37. Powdered iron is generally used as the core in a high-frequency coil.
38. Skin effect refers to the situation where at high frequencies current tends to flow in the outer part of a conductor.
39. An impedance match assures that maximum power is transferred from the source to the load.
40. The primary and secondary impedances are related by the square of the turns ratio.
41. 15.8 to 1 — The required turns ratio is $\sqrt{1000/4} = \sqrt{250} = 15.8$.



1091A

1. C - AN ADVANTAGE OF A PUSH-PULL AMPLIFIER AS COMPARED TO A SINGLE-ENDED STAGE IS -- CANCELLATION OF EVEN HARMONIC DISTORTION.
The push-pull stage cancels even harmonic distortion and provides increased power output.
2. C - CROSSOVER DISTORTION IN A PUSH-PULL CLASS B AMPLIFIER IS REDUCED BY -- PROVIDING A SMALL AMOUNT OF FORWARD BIAS.
The forward bias aligns the transfer curves so that the distortion near zero bias is eliminated.
3. B - A COMPLEMENTARY-SYMMETRY AMPLIFIER USES -- A PNP AND AN NPN TRANSISTOR.
The complementary-symmetry amplifier uses a PNP and NPN transistor, thus eliminating the need for a phase inverter circuit.
4. C - RC COUPLING IS UNDESIRABLE IN A -- CLASS B AMPLIFIER.
RC coupling is undesirable in a Class B amplifier since base current can develop reverse bias on the coupling capacitor.
5. C - FOR GOOD LOW-FREQUENCY RESPONSE IN AN RC COUPLED AMPLIFIER, -- THE COUPLING CIRCUIT TIME CONSTANT SHOULD BE LONG.
With a long time constant, little signal voltage is developed across the coupling capacitor.
6. B - FOR GOOD HIGH-FREQUENCY RESPONSE IN AN RC COUPLED AMPLIFIER, -- THE SHUNT CAPACITY SHOULD BE LOW.
The high-frequency response is controlled by the value of shunt capacity.
7. D - IN A TUNED H-F AMPLIFIER A HIGH CIRCUIT Q PRODUCES -- HIGH GAIN AND A NARROW BANDWIDTH.
A high circuit Q is obtained with a high coil Q (low series R) and a high value of shunt resistance.
8. A - WHEN THE COUPLING IN A DOUBLE-TUNED AMPLIFIER IS INCREASED BEYOND THE CRITICAL VALUE (OVERCOUPLED), -- BANDWIDTH INCREASES.
Overcoupling develops a double humped response curve increasing bandwidth.
9. D - THE COMPONENTS IN FIGURE 22 THAT ELIMINATE THE EFFECTS OF FEEDBACK ARE -- R_N AND C_N .
 R_N and C_N form the neutralization circuit.
10. D - FOR MAXIMUM EFFICIENCY, A H-F POWER AMPLIFIER IS OPERATED -- CLASS C.
Maximum efficiency is obtained with Class C operation.

PRACTICE EXERCISE SOLUTIONS (Continued)

31. Connecting the base of Q_2 to the top of L_2 would shunt the tank circuit with a low impedance. This would reduce circuit Q , reducing gain and increasing bandwidth.
32. The interelectrode capacity of a transistor becomes important at high frequencies where its reactance decreases.
33. False — Due to the center tapped arrangement of T_2 , the signal at the bottom is 180° out-of-phase with the signal at the top.
34. Push-pull operation provides increased power output and cancellation of even harmonic distortion.
35. Zero bias provides Class B operation.
36. A power amplifier is operated Class B or C to obtain high efficiency.
37. Powdered iron is generally used as the core in a high-frequency coil.
38. Skin effect refers to the situation where at high frequencies current tends to flow in the outer part of a conductor.
39. An impedance match assures that maximum power is transferred from the source to the load.
40. The primary and secondary impedances are related by the square of the turns ratio.
41. 15.8 to 1 — The required turns ratio is $\sqrt{1000/4} = \sqrt{250} = 15.8$.

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1091A

Example:

A	☐
B	☐
C	☐
D	☐

Automobiles usually run on

A. four wheels. B. two wheels. C. five wheels. D. no wheels.

1.

A	☐
B	☐
C	☒
D	☐

An advantage of a push-pull amplifier as compared to a single-ended stage is
(A) cancellation of odd harmonic distortion. (B) reduced power output. (C) cancellation of even harmonic distortion. (D) reduced output voltage swing.

2.

A	☐
B	☐
C	☒
D	☐

Crossover distortion in a push-pull Class B amplifier is reduced by
(A) operating Class C. (B) providing a small amount of reverse bias. (C) providing a small amount of forward bias. (D) decreasing the input signal amplitude.

3.

A	☐
B	☒
C	☐
D	☐

A complementary-symmetry circuit uses
(A) two NPN transistors. (B) a PNP and an NPN transistor. (C) two PNP transistors. (D) a phase inverter.

4.

A	☐
B	☐
C	☒
D	☐

RC coupling is undesirable in a
(A) low-frequency amplifier. (B) Class A amplifier. (C) Class B amplifier. (D) audio amplifier.

5.

A	☐
B	☐
C	☒
D	☐

For good low-frequency response in an RC coupled amplifier,
(A) the coupling capacitor should have a low value. (B) the coupling circuit time constant should be short. (C) the coupling circuit time constant should be long. (D) the value of shunt capacity should be high.

6.

A	☐
B	☒
C	☐
D	☐

For good high-frequency response in an RC coupled amplifier,
(A) the shunt capacity should be high. (B) the shunt capacity should be low. (C) the coupling capacitor should have a very small value. (D) the coupling circuit time constant should be short.

7.

A	☐
B	☐
C	☐
D	☒

In a tuned h-f amplifier a high circuit Q produces
(A) high gain and a wide bandwidth. (B) low gain and a wide bandwidth. (C) low gain and a narrow bandwidth. (D) high gain and a narrow bandwidth.

8.

A	☒
B	☐
C	☐
D	☐

When the coupling in a double-tuned amplifier is increased beyond the critical value (overcoupled),
(A) bandwidth increases. (B) gain increases. (C) bandwidth decreases. (D) gain drops to zero.

9.

A	☐
B	☐
C	☐
D	☒

The components in Figure 22 that eliminate the effects of feedback are
(A) R'_b and C_{BC} . (B) R_2 and C_3 . (C) C_2 . (D) R_N and C_N .

10.

A	☐
B	☐
C	☐
D	☒

For maximum efficiency, a h-f power amplifier is operated
(A) Class A. (B) Class B. (C) Class AB. (D) Class C.



VACUUM TUBE OSCILLATORS

1094

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





An oscillator operating in the range of 50 to 100 kHz is used in a tape recorder to erase the tape as well as provide "bias" to reduce distortion during recording.

Courtesy Crown International

CONTENTS

Oscillators and Their Applications	Page 5
How Electric Oscillations Are Developed	Page 5
An Oscillating Circuit	Page 7
Fundamental Oscillator Action	Page 9
Grid Leak Biased Oscillator	Page 12
Types of Oscillators	Page 13
The Tuned-plate Oscillator	Page 14
The Tuned-grid Oscillator	Page 15
The Hartley Oscillator	Page 15
The Colpitts Oscillator	Page 18
The Tuned-plate Tuned-grid Oscillator	Page 19
The Electron-Coupled Oscillator	Page 20
The Clapp Oscillator	Page 22
The Crystal Oscillator	Page 23
Crystal Cuts	Page 23
Temperature Coefficient of Crystals	Page 24
Crystal Oscillator Circuits	Page 24
Parasitic Oscillations	Page 27
RC Oscillators	Page 28
The Multivibrator	Page 28
The Cathode-Coupled Multivibrator	Page 31
The Phase Shift Oscillator	Page 32
The Wein Bridge Oscillator	Page 34
The Blocking Oscillator	Page 35

*The greater a man's knowledge of what has been done, the
greater is his power of knowing what to do.*

—Selected

VACUUM TUBE OSCILLATORS

OSCILLATORS are one of the most widely used type of circuit in the broad field of electronics. Oscillators are used to generate the high-frequency signals used in radio, television, communications, and radar transmitters and receivers. They are also used in industry to time and synchronize many automatic control processes, heat and melt various metals, seal plastic products and as a source of signals used in measuring the thickness of various materials.

Ultrasonic oscillators have many uses such as; cleaning various parts, homogenizing and mixing liquids, and machining delicate parts. Many of the instruments used by the medical profession, such as diathermy machines, contain electronic oscillators. High-frequency oscillators are also extensively used in the preparation of food in such devices as electronic ovens, sometimes called "radar" ovens. To list all the possible uses of oscillators would take many pages. It is therefore important that you be familiar with the operation of basic oscillator circuits.

According to the dictionary, the word "oscillate" means to vibrate, to swing back and forth, or to pass from one state to another and back again. This is a rather broad definition as no mention is made regarding the nature of the variations, which may be gradual or abrupt. Also, no restrictions are placed on the rate of variation. These might occur once a century or millions of times a second. Thus, it can be stated that the weather oscillates each year because it passes from winter cold to summer heat and back again to winter cold.

In mechanical systems, oscillation refers to a back and forth motion around a center, or fixed point, to produce movements like those of a clock pendulum, or an automobile windshield wiper.

In electric power systems, the periodic reversals of alternating current and voltage also are oscillations. Each complete series of changes constitutes a cycle, and the number of cycles which occur in a second determines the frequency, measured in Hertz. Thus, the common electric power circuits are rated simply as "110 - 120 volt, 60 Hz ac" and it is understood that the voltage and current oscillate 60 times per second.

In radio, television, and other electronic applications, the term "oscillate" has a special meaning. It refers almost entirely to the generation of ac by means of electron tubes and transistors. There are no definite limits to the

frequencies of these oscillations, which range from a few Hz to several thousand million Hz. In this lesson we shall examine the methods used to generate ac signals with vacuum tubes.

OSCILLATORS AND THEIR APPLICATIONS

The primary purpose of an electron tube **OSCILLATOR** is to convert direct current energy into alternating current energy. AMPLIFIER tubes are controlled by an external signal source connected into the grid circuit and change the dc power from the plate supply into ac power in the plate load.

Although the *basic action* of the tube does not change, when used as an oscillator it functions without the aid of an external source of signal voltage. Thus, the electron tube oscillator is essentially an alternating current generator. The waveform, amplitude and frequency of the oscillations which it generates depend only upon the circuit elements used and the manner in which they are connected, rather than on some external signal injected into the grid circuit.

HOW ELECTRIC OSCILLATIONS ARE DEVELOPED

The essential components of an electron tube oscillator are: (1) a high voltage dc source which furnishes the direct current power; (2) a tuned circuit, consisting of capacitive and inductive elements, in which the oscillations are produced; and (3) a switch, or its equivalent, that controls the period during which dc power is supplied to the tuned circuit.

The actions which produce electronic oscillations can be explained by means of the simple circuit illustrated in Figure 1. Inductor L and capacitor C are connected in parallel to form a tank circuit. To simplify the analysis, we shall assume a negligible resistance in the circuit.

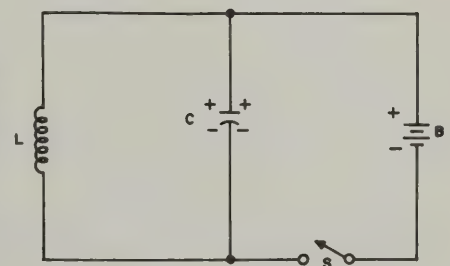


Figure 1

With the aid of switch S , the tank circuit may be connected momentarily to the dc supply B . Closing the switch for an instant permits capacitor C to charge to the supply voltage. This makes the lower plate of C negative and its upper plate positive. Due to the almost zero resistance of the charging circuit, the charge occurs very rapidly. At the same time, the current in the inductor begins to increase. However, the self induction caused by the expanding magnetic field prevents an appreciable increase in the current in this short duration of time. Thus, a definite quantity of electric energy is stored in the charged capacitor but very little is in the inductor.

At the instant the switch is opened, the circuit conditions are as illustrated in Figure 2A. Capacitor C is in an unnatural state. Like a spring that has been stretched out of shape, it tries to return to its normal state. Consequently, capacitor C discharges through inductor L, with the electrons flowing from the negative plate up through the inductor to the positive plate. This action continues until the capacitor is completely discharged as shown in Figure 2B.

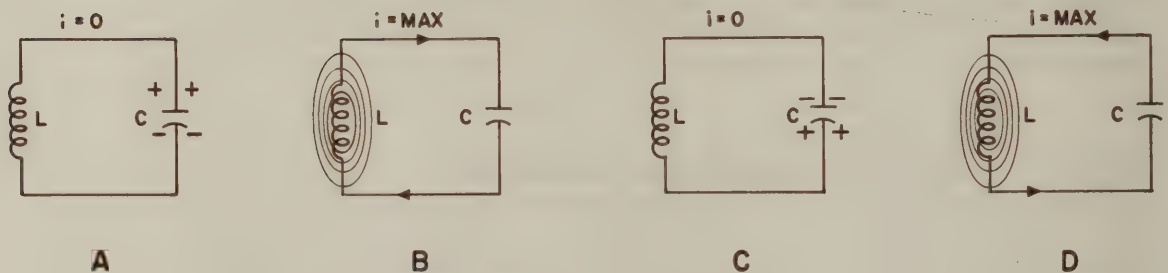


Figure 2

The action does not stop at this point, however. The electron flow through the inductor produces a magnetic field around the inductor as pictured in Figure 2B.

When the capacitor is discharged, the magnetic field begins to collapse. The resulting induced emf makes the upper end of the inductor negative with respect to the lower end. Now the inductor functions as the voltage source, and capacitor C charges to the polarity shown in Figure 2C.

Once more, the capacitor is in an unnatural state. Therefore, it discharges through L and produces a magnetic field around the inductor as shown in Figure 2D. This time, the magnetic field is opposite to that of Figure 2B. When the field collapses, the induced emf makes the lower end of the inductor negative with respect to the upper end. The resulting current charges capacitor C to the polarity shown in Figure 2A. With the original conditions restored, one cycle has been completed. As before, the action does not cease at this point. Since we have assumed a negligible resistance in the circuit, little energy is lost by the circuit and the cycle of events repeats over and over until circuit losses eventually absorb all the energy.

Thus, a periodic transfer of energy from the electrostatic charge of the capacitor to the electromagnetic field of the inductor and back to the capacitor produces an oscillating current in the circuit. The rate at which these trans-

fers occur depends upon the capacitance and the inductance in the tank circuit.

For example, IF THE CAPACITANCE IS INCREASED, more time is required for its discharge and charge; therefore, the FREQUENCY OF THE OSCILLATIONS IS REDUCED. An increase of inductance produces the same result because of the larger counter emf developed for the same change in current. On the other hand, a DECREASE OF EITHER INDUCTANCE OR CAPACITANCE CAUSES AN INCREASE IN THE OSCILLATION FREQUENCY because less energy is stored in the capacitor or a smaller counter emf opposes the current changes.

AN OSCILLATING CIRCUIT

If a condition of zero resistance is assumed, the instantaneous voltages and currents for the complete cycle, as explained for Figure 2, can be shown by the graph of Figure 3. The instantaneous circuit current is represented by curve "i", the capacitor voltage by curve " e_c ", and the voltage across the inductor by curve " e_L ".

During oscillation the capacitor and inductor alternately serve as sources of energy and as loads. These facts are emphasized by drawing the voltage curves with solid and dashed lines. The solid portions of the curves represent the time when the component is acting as a source, and the dashed portions show when the component is a load. Since the capacitor and inductor are in series at all times, the curves show that the voltages across L and C are always equal but opposite in phase.

The time scale of Figure 3 is lettered to correspond with the conditions of Figure 2. Starting at Point A, which corresponds to Figure 2A, the capacitor is charged to maximum voltage, the current is zero, and the voltage applied across the inductor is equal to the capacitor voltage. As the capacitor discharges, its voltage decreases as the current increases. This action continues until at Point B the current is maximum and the capacitor voltage is zero.

During the time interval between A and B, the capacitor is the voltage source and the inductor is the load. As explained for Figure 2B, the increase of current causes an expanding field around the inductor, which by self induction induces a counter emf. Since it was assumed that this simple circuit has negligible resistance, this action causes the current to lag the voltage by 90° . Thus, at the instant the capacitor voltage is zero, the peak current causes a maximum magnetic field around the coil.

Immediately after reaching its maximum, the current starts to decrease, thereby causing the magnetic field to contract. This action induces a voltage of the correct polarity to maintain the current. Because the induced emf is proportional to the rate at which the current changes, and because the current changes most rapidly as it passes through zero, the voltage induced in the inductor is maximum at the zero current point as shown at Point C in Figure 3. While inductor L is acting as a source of voltage, the curves show that capacitor C is acting as a load and is being charged, this time with the top plate becoming negative.

With zero current at Point C, there is no magnetic field around the coil; therefore, no electromagnetic induction takes place. The capacitor has been charged to the maximum voltage, the increase of which is indicated by the broken line section of curve e_c , between Points B and C. During this interval, the self induced emf in the inductor is the source voltage and the capacitor is the load. Therefore, when the capacitor is being charged, the current leads the capacitor voltage by 90° .

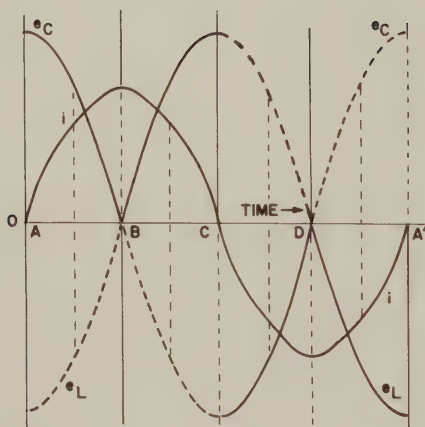
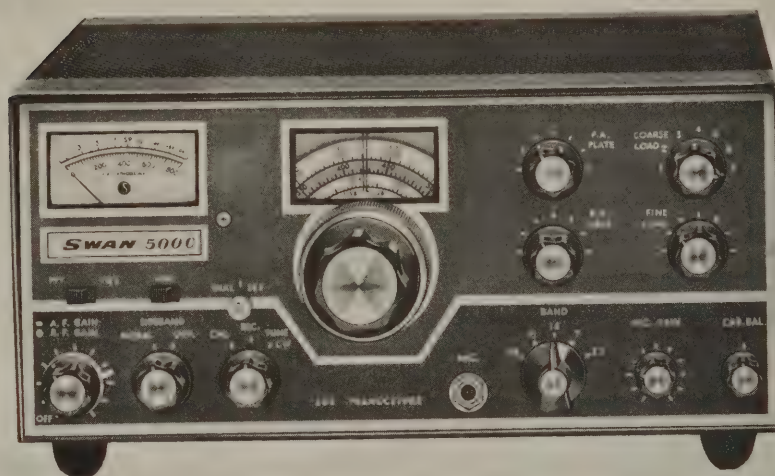


Figure
3



Both fixed and variable frequency LC oscillators are employed in communications transceivers like the one shown above.

Courtesy Swan Electronics Corp.

Except for the reversal of polarity, the various values at Point C are the same as those of Point A; therefore, with this exception, the actions through Points D and A' are the same as those explained through Points A, B, and C.

The complete cycle consists of four distinct actions. During the interval A-B, the capacitor discharges and transfers its stored energy to the magnetic

field around the coil. During the interval B-C, the magnetic field collapses and by electromagnetic induction transfers its energy to the capacitor, charging it to the opposite polarity. During the interval C-D, the capacitor discharges, and transfers its stored energy to the magnetic field, which transfers the energy back to the capacitor during interval D-A'.

Notice that the capacitor charges and discharges during each alternation of current and, therefore, it is charged and discharged TWICE for each cycle.

The greater the inductance, the longer the time required for capacitor C to discharge through it. Likewise, the greater the capacitance, the longer it will take to charge or discharge. Thus, the inductance and capacitance both determine the period of time required for one complete cycle or oscillation.

The fundamental oscillation frequency of any LC circuit is the same as its resonant frequency, which as you know is:

$$f = \frac{1}{2\pi\sqrt{LC}} \text{ or } \frac{.159}{\sqrt{LC}}$$

FUNDAMENTAL OSCILLATOR ACTION

When a pulse of electric energy is applied to the LC circuit of Figure 1, the oscillations are initiated, and would continue indefinitely if no resistance was present in the circuit. However, resistance is present and the tank circuit actually contains inductance, capacitance and resistance, as indicated in Figure 4. The resistance consists of the resistance of the inductor, the connecting wires, and the capacitor plates. Resistance produces heat and, because of the resulting loss, some energy is dissipated during each cycle. Unless this energy is replaced by an external source, the capacitor is charged to a slightly lower voltage during each successive alternation. A graph of the capacitor voltage appears as shown by the decreasing amplitude or "damped wave" curve of Figure 5.

The number of cycles which occur before the oscillations die out completely is inversely proportional to the resistance in the circuit. Hence, to maintain oscillations, the LC circuit must be supplied regularly with additional energy to replace that lost by the circuit resistance.

As indicated by the curve of Figure 5, the oscillatory circuit has the property of storing and releasing more energy than it loses per cycle. Therefore, it is not necessary to supply energy continuously to maintain the oscillations.

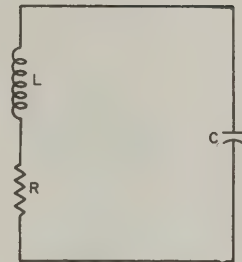


Figure
4

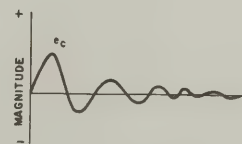


Figure
5

Instead, a switching circuit can be arranged to connect the external source to the LC circuit only during a small portion of each cycle. Therefore, enough energy to make up for that lost during the remainder of the cycle is supplied to the tank circuit.

At the relatively high frequencies generated by most common oscillators, it is impractical to attempt to use any mechanical device to serve as the switch, as in Figure 1. Instead, an electron tube circuit which can switch at a high speed is used. A simplified version of an electron tube oscillator circuit is shown in Figure 6. It includes the oscillatory LC tank circuit consisting of L_2 and C_2 .

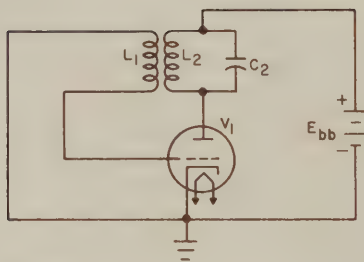


Figure
6

The plate circuit of triode tube V_1 is connected in series between the tank circuit and the power supply negative terminal. Since the grid voltage of a triode tube controls the plate current, it also controls the dc plate resistance of the tube. When the grid becomes more positive with respect to the cathode, the internal resistance decreases and plate current increases. When the grid becomes more negative, the internal resistance increases and plate current decreases.

When power is first applied to the circuit, plate current increases towards its saturated value since there is no bias on V_1 . This causes a surge of current in the tank circuit L_2C_2 . This surge charges capacitor C_2 and, at the same time, sets up an expanding magnetic field around inductor L_2 that cuts coil L_1 and induces a voltage in it. This induced voltage is of a polarity which makes the grid positive with respect to the cathode, thereby increasing the conductivity and plate current.

As this action continues, the plate current increases and the internal resistance of the tube decreases to allow almost full plate supply voltage across the coil and capacitor. Thus, the L_2C_2 circuit condition corresponds to that shown in Figure 1 when the switch is closed. At this instant, the capacitor is fully charged; therefore, the current into it drops to zero. Neglecting the V_1 plate current and considering only the oscillation or tank current, the circuit condition at this time corresponds to that shown by the curve at Point A in Figure 3.

Going back to the circuit of Figure 6, as plate current approaches maximum, the induced voltage in coil L_1 dies out. Therefore, the grid becomes less positive with respect to the cathode. This reduces the plate current and lowers the voltage applied to the L_2C_2 tank. As a result, C_2 discharges through inductor L_2 and starts the oscillations as explained for Figure 2. When the total coil current decreases, there is a contraction of flux lines around L_2 and

the resulting voltage induced in L_1 is the opposite polarity which drives the grid of V_1 negative.

When the grid is driven negative to the plate current cutoff point, the tube becomes nonconductive and the tank circuit conditions of Figure 6 duplicate those of Figure 1 with the switch open. The oscillating current and the induced voltage in L_1 maintain the grid at plate current cutoff as the cycle continues. Finally, as C_2 is just about discharged, the negative voltage induced in L_1 decreases to less than cutoff, the tube becomes conductive, and the cycle repeats.

Thus, the electron tube is conductive during one part of each cycle and non-conductive during the remainder of each cycle of the oscillating current. When conductive, the tube allows the capacitor to charge to the supply voltage and thus replaces the energy lost during the preceding cycle. As a result, the oscillations continue as long as the external power is applied to the circuit.

In the circuit of Figure 6, the tube plate current varies from maximum to cutoff. Larger changes of grid voltage do not increase these variations, hence the amplitude of the oscillations is determined by the plate current-grid voltage characteristics of the tube. For this reason, the dc energy supplied to the oscillatory circuit is limited by the tube, which tends to maintain the amplitude of the oscillations at a fixed level as pictured in Figure 7.

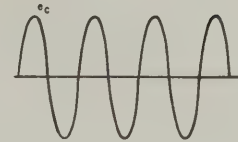


Figure
7

GRID LEAK BIASED OSCILLATOR

As previously mentioned, an LC or tank circuit is capable of storing much more energy than is lost during any one cycle of oscillation. Therefore, a short pulse of dc from the power supply is sufficient to maintain oscillation.

In the simplified circuit of Figure 6, the tube is conductive for a considerable part of a cycle. Therefore, it dissipates more energy than is actually required. As a result, the operation is not very efficient and the plate current may be of sufficient amplitude to overheat and damage the tube.

To provide a more efficient and practical oscillator, the tube is generally operated Class C with a bias of about twice that required for plate current cutoff. A pulse of plate current sufficient to maintain oscillation then occurs only during the most positive peaks of the grid voltage swing.

If a fixed bias arrangement is employed, the negative grid voltage prevents the initial pulse of current when power is applied. As previously explained, the oscillations cannot start unless the tank circuit is supplied with energy by an initial pulse of plate current. Hence, to be self-starting, an oscillator of this type requires an initial zero bias which builds up to the normal value after the circuit is in operation.

This type of bias voltage is provided by employing the grid leak bias arrangement shown in the circuit of Figure 8. The circuit contains all of the components of Figure 6, with the addition of the R_1C_1 bias circuit and bypass capacitor C_3 . Capacitor C_3 provides a low reactance path for the high frequency ac between the + E_{bb} end of the tank circuit and the cathode of the tube.

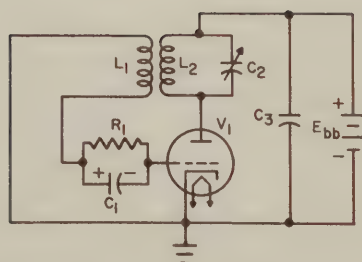


Figure
8

Before the plate supply circuit is closed, the oscillator of Figure 8 is electrically at rest, and there is zero bias on the grid of tube V_1 . When power is applied, a pulse of plate current occurs, and the L_2C_2 circuit goes into oscillation as explained for Figure 6.

Whenever the grid is positive with respect to the cathode, there is grid current which causes a voltage drop across resistor R_1 and charges capacitor C_1 to the indicated polarity. Usually the resistance of inductor L_1 is very low, and the dc drop across it is negligible. Thus, the positive plate of C_1 can be considered as being at dc ground potential. Since the cathode of tube V_1 is grounded also, the voltage across R_1C_1 is applied to the grid-cathode circuit and the grid becomes negative with respect to the cathode.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

OSCILLATORS AND THEIR APPLICATIONS

HOW ELECTRIC OSCILLATIONS ARE DEVELOPED

AN OSCILLATING CIRCUIT

FUNDAMENTAL OSCILLATOR ACTION

1. What is the function of an oscillator?
2. Does an oscillator require a signal source to generate an ac signal?
3. How does a parallel LC or "tank" circuit develop an alternating current signal?
4. What effect does the value of inductance and capacitance have on the frequency of the ac signal developed by a tank circuit?
5. As shown in Figure 3, the capacitor is serving as the source during intervals B-C and D-A'. True or False?
6. What is the frequency of oscillation of a parallel LC circuit if the value of C is 100 pF and the value of L is 400 μ H?
7. What is the name given to the waveform shown in Figure 5?
8. What factor in an LC circuit makes it necessary to continually supply it with energy to develop continuous oscillations?
9. L_1 and L_2 in Figure 6 must be phased so that as plate current is increasing, the grid of V_1 is driven (positive) (negative)?
10. What happens to the grid voltage on V_1 in Figure 6 as C_2 begins to discharge?
11. If V_1 in Figure 6 were biased at cutoff before power were applied, would oscillations occur when power was applied?

The grid is driven positive only during small portions of the oscillation cycle. Therefore, the grid current occurs in pulses, each of which charges capacitor C_1 as previously explained. Between the pulses of grid current, capacitor C_1 discharges through resistor R_1 .

In practice, the values chosen for C_1 and R_1 are relatively large for a given oscillation frequency. Hence, the large time constant permits only a slight discharge of the R_1C_1 network before it is recharged by a succeeding pulse of grid current. Thus, as the oscillations attain their normal operating amplitude, an essentially constant grid bias voltage is developed across C_1 and parallel connected resistor R_1 .

Although the primary purpose of the R_1C_1 bias arrangement of Figure 8 is to provide more efficient operation of the tube, the action permits self-starting and serves as an automatic amplitude control. The positive grid swing and the charge on capacitor C_1 depend upon the amplitude of the voltage induced in coil L_1 . Therefore, any increase in the amplitude of the oscillations in the tank circuit results in a greater induced voltage which increases the grid current. The increased grid current causes a greater drop across resistor R_1 and, thereby, increases the grid bias.

With greater negative grid bias, the plate current pulses are of shorter duration, and less energy is supplied to the tank circuit. Thus, the complete action tends to reduce the oscillation amplitude. On the other hand, a decrease in the amplitude of the oscillations results in a lower negative grid bias. This allows more energy to be supplied to the oscillatory tank, and thus tends to increase the amplitude.

TYPES OF OSCILLATORS

Electron tube oscillators may be classified in a number of ways, such as:

(1) according to the frequency of the generated ac, (2) according to the waveform of the output voltage, and (3) with regard to the circuit arrangement which makes oscillation possible.

A sharp line cannot be drawn between the various types included in the first two classes. For example, an oscillator which produces radio frequencies in the broadcast range is a high-frequency oscillator compared to an audio-frequency signal generator, but is a low-frequency device compared to a "microwave" oscillator.

Waveform classification places oscillators which generate pure or nearly pure sinewaves into a group called SINUSOIDAL. All others are classed as NON-

SINUSOIDAL. This grouping is made in accordance with the fact that different methods of analysis may be applied to oscillators, depending upon whether their output is sinusoidal or not.

However, again the distinction between types is not sharp. In many oscillators the output wave can be varied without interruption from sinusoidal to non-sinusoidal by adjusting some circuit element. The circuit of Figure 8 is an example of a sinewave or sinusoidal oscillator. The multivibrator which is taken up later in this lesson is an example of the nonsinusoidal type.

The method of classifying oscillators in accordance with the circuit arrangements used to produce the oscillations permits a sharp distinction to be made between oscillator circuits.

Oscillations can be set up and sustained in a tank circuit if means are provided to replace the energy lost in the circuit resistance. This is necessary because resistance absorbs power from a circuit and dissipates it in the form of heat. Therefore, an oscillator requires some element which can supply power to replace that which is lost.

In the circuits of Figures 6 and 8, the energy lost in the tank circuit is replaced from the power supply each time the tube conducts. Thus, sustained oscillations are maintained. Due to the inductive coupling between coils L_1 and L_2 , a part of the energy in tuned tank circuit L_2C_2 is fed into the grid circuit with proper phase and magnitude to cause the tube to conduct at just the right instant. Hence, AN OSCILLATOR IS A SELF-EXCITED AMPLIFIER IN WHICH A PORTION OF THE OUTPUT ENERGY IS FED BACK TO THE INPUT CIRCUIT. This idea is pictured in the diagram of Figure 9 where the input voltage is increased in amplitude by means of the normal amplifying properties of the tube V_1 , and the output voltage is fed back to the input, in proper phase and amplitude, through some type of coupling network. The feedback must be positive since the output must aid the input. The feedback amplitude must be sufficient to make up for circuit losses.

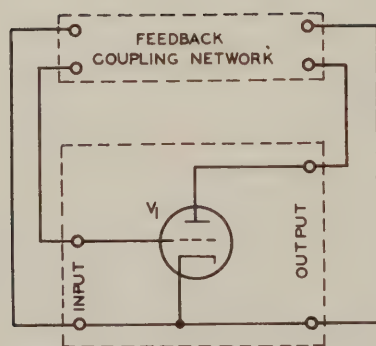


Figure
9

Since the output contains more ac energy than needed to drive the input grid of the amplifier tube properly, the surplus power is available for external use. It may be coupled to the input of a conventional amplifier or other circuit. Though differing in details of circuit arrangement and components, all of the feedback oscillators to be described are of this basic form.

THE TUNED-PLATE OSCILLATOR

The circuit of Figure 8 satisfies the requirements of a feedback oscillator because a change in the grid voltage of V_1 is amplified by the tube and ap-

pears in the plate circuit. By means of the inductive coupling between coils L_1 and L_2 , energy from the plate circuit is fed back to the grid and amplified again.

Connections between coil L_1 and the grid circuit are made so that the grid is driven positive and the conductivity of the tube is increased as the plate current increases. The plate voltage decreases as the plate current increases. Thus, the tube provides the power necessary to sustain oscillations. The tuned tank circuit which determines the frequency of oscillation is in the plate circuit of Figure 8. Therefore, this arrangement is known as a **TUNED-PLATE OSCILLATOR**.

THE TUNED-GRID OSCILLATOR

The tuned circuit may also be located in the grid circuit of the tube, with the feedback coupling coil in the plate circuit. This arrangement, called a **TUNED-GRID OSCILLATOR** or **ARMSTRONG OSCILLATOR**, is shown schematically in Figure 10. Any change of plate current in inductor L_1 , sometimes called a **TICKLER COIL**, induces a corresponding voltage in inductor L_2 . As coil L_2 is connected across the grid circuit, energy is fed back from the plate circuit to the grid by the mutual inductance between L_1 and L_2 . As in the circuit of Figure 8, the coupling between the coils must be such that the feedback voltage has sufficient amplitude and proper phase to sustain oscillations.

Compared with Figure 8, the circuit of Figure 10 shows another arrangement of grid-leak bias components. In this case the grid leak components are R_1 and C_2 . So far as the action in developing the grid bias is concerned, the two arrangements function alike. However, with resistor R_1 in parallel with the tuned circuit, as in Figure 10, it also acts as a "load" for the tank circuit and thus provides greater stability of oscillation. Therefore, this method offers a definite advantage and is employed in most oscillator circuits of this type. The grid leak bias arrangement of Figure 10 is referred to as **SHUNT GRID LEAK BIAS** while that used in Figure 8 is referred to as **SERIES GRID LEAK BIAS**.

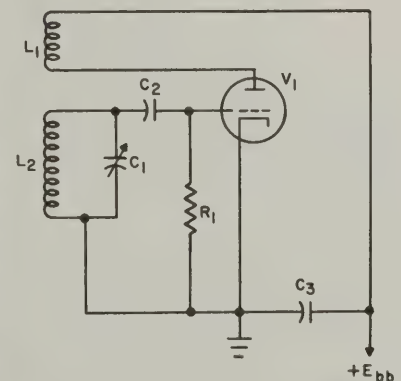


Figure 10

THE HARTLEY OSCILLATOR

Perhaps the most common oscillator circuit is the **HARTLEY OSCILLATOR**, one form of which is shown in Figure 11. Here, tapped tank coil L_1 , made up of sections L_A and L_B , performs the functions of coils L_1 and L_2 , respectively of Figure 10.

Tracing the dc circuits, there is a path from the cathode through the tube to the plate, through coil L_A and the power supply from $+E_{bb}$ to ground and back to the cathode. The grid connects to the cathode through resistor R_1 .

Tracing the ac circuits, there is a path from cathode to grid, through capacitor C_2 , coil L_B and capacitor C_4 back to the cathode. Another ac path is from cathode to plate, through coil L_A , and capacitor C_4 back to the cathode.

The reactance of a capacitor is extremely high at dc or zero frequency and decreases as the frequency increases. Thus, for the dc paths, C_2 and C_4 act as open circuits and are known as "blocking capacitors". For the comparatively high oscillation frequency, the reactance of these capacitors is so low that the tap between L_A and L_B is at approximately ac ground potential and the lower end of L_B connects to the tube grid. Thus, the ac voltages produced in section L_B of the coil are impressed across R_1 and the grid-cathode circuit with negligible loss.

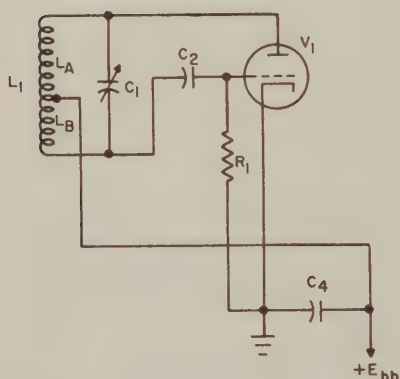


Figure 11

The operation of the Hartley oscillator of Figure 11 is essentially the same as that of the tuned-grid oscillator of Figure 10, but with one important difference. In the tuned-grid oscillator of Figure 10, the tank circuit consists of inductor L_2 and capacitor C_1 while the plate energy is fed back through the inductive coupling between coils L_1 and L_2 . In the Hartley oscillator of Figure 11, the tank circuit consists of coil sections L_A and L_B . Connected in series across capacitor C_1 , both carry the oscillating tank current. Thus, even when no inductive coupling is provided between the sections of the coil, plate energy is fed back to the grid and the circuit oscillates.

For convenience and economy, sections L_A and L_B are wound as a single coil and proper phase relations are obtained by connecting the plate to one end and coupling the grid circuit to the other. The amplitude of the feedback voltage is determined by the position of the tap.

In Figure 11, the plate supply is in series with inductor L_A and the tube plate circuit. Therefore, the dc plate current is carried by part of the tank coil. When the direct current of the tube plate circuit passes through a portion of the plate load, the arrangement is called a **SERIES FEED**. Hence, the circuit in Figure 11 is known as a **SERIES-FED HARTLEY OSCILLATOR**.

A disadvantage of the series fed arrangement in Figure 11 is that the tank circuit is placed at the high dc potential of the power supply with respect to ground. This constitutes a shock hazard to the operator of large equipment when tank adjustments have to be made. This danger of shock can be eliminated by rearranging the circuit so as to remove the high direct voltage from

the tank circuit. One method of doing this is pictured in Figure 12. The tank inductor tap is connected directly to the V_1 cathode and ground. Capacitor C_3 is connected between the tank and the tube plate to complete the feedback path.

A power supply normally has a very low impedance between its terminals and short circuits any ac voltage on the V_1 plate when the $+E_{bb}$ terminal is connected directly to the tube plate. To prevent this condition, a resistor or radio-frequency choke (RFC) is included between the tube plate and the E_{bb} circuit as shown in Figure 12. The resistor or choke presents a high impedance to the ac energy and prevents the power supply from short circuiting it.

On the other hand, capacitor C_3 presents a low impedance to the ac voltage and couples it to the tank circuit to sustain oscillation. In addition, capacitor C_3 blocks the dc and prevents the power supply from shorting to ground through RFC and L_A .

When the plate current is carried by a path that is parallel to and separate from the plate load, the arrangement is known as SHUNT FEED. Hence, the circuit of Figure 12 is that of a SHUNT-FED HARTLEY OSCILLATOR.

Although shown here for the plate circuit of the Hartley oscillator, series and shunt feed may be used in the plate circuits of most electron tube oscillators. In each case, **THE CIRCUIT IS SERIES-FED IF THE ELECTRODE DIRECT CURRENT PASSES THROUGH ANY PART OF THE LOAD, AND IS SHUNT-FED IF IT DOES NOT.**

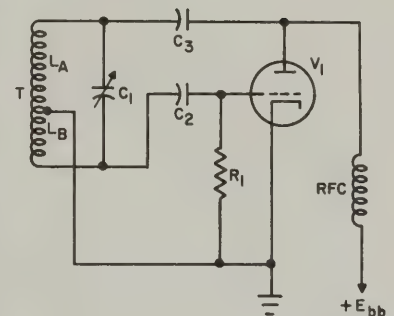


Figure 12

THE COLPITTS OSCILLATOR

From a dc standpoint, a capacitor is nonconductive and prevents current through it. However, when connected across an ac supply, a capacitor periodically charges and discharges and the resulting displacement currents appear as alternating current in its circuit. The magnitude of this ac depends upon the applied voltage and the capacitive reactance, which is measured in ohms the same as resistance. Thus, in ac and especially r-f circuits, the displacement currents into and out of a capacitor are treated the same as the currents through a resistor.

In the Hartley circuit of Figure 12, the coil is tapped to provide the feedback voltage, but, to obtain the same action in the **COLPITTS OSCILLATOR** circuit of Figure 13A, the tank capacitor is tapped. In other respects, the circuits are alike and operate on the same principles.

From a practical standpoint, the tank capacitor is divided into two parts, C_A and C_B , with the grid circuit across C_B , the plate circuit across C_A , and coil L across both. The feedback is controlled by changing the capacitance of C_A or C_B . The smaller the capacitance of C_B , the higher its reactance and the greater the voltage developed across it.

Since the tube cathode connects to the junction between the parts of the tank capacitor, neither of which carries dc, all Colpitts oscillators must be shunt-fed.

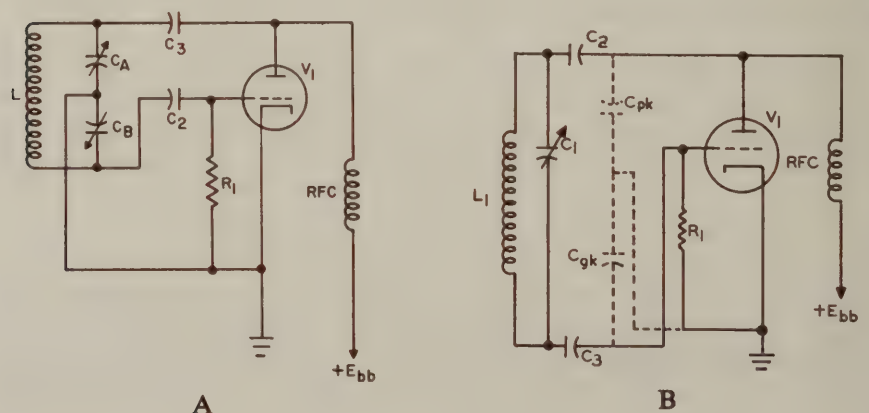


Figure 13

While they are quite similar in operation, both the Hartley and Colpitts oscillators have their own advantages. For example, the Hartley is simpler

The following Practice Exercise questions cover the subjects which you have just studied. They are:

GRID LEAK BIASED OSCILLATOR

TYPES OF OSCILLATORS

THE TUNED-PLATE OSCILLATOR

THE TUNED-GRID OSCILLATOR

THE HARTLEY OSCILLATOR

12. When power is first applied to the circuit of Figure 8, what is the value of grid bias on V_1 ?
13. Explain briefly how bias is developed in the circuit of Figure 8.
14. For proper operation in the circuit of Figure 8, R_1 and C_1 should have a short time constant. True or False?
15. Assume that the amplitude of the oscillations in the circuit of Figure 8 decreases. Explain how the grid leak bias arrangement helps to counteract the decrease.
16. How would you describe an oscillator circuit in terms of an amplifier?
17. Is the feedback employed in an oscillator circuit positive or negative?
18. What is the major difference between the circuits of Figures 8 and 10?
19. What is the name given to L_1 in the circuit of Figure 10?
20. How is the feedback obtained in the tuned-grid and tuned-plate oscillators?
21. How is feedback obtained in a Hartley oscillator?
22. What is the function of C_2 and R_1 in the circuit of Figure 11?
23. Suppose the top and the bottom ends of L_B in Figure 11 are reversed. Would the circuit still oscillate?
24. What is the difference between a series-fed oscillator and a shunt-fed oscillator?
25. What is the disadvantage of a series-fed oscillator arrangement?
26. What is the purpose of RFC in the circuit of Figure 12?

to tune, but the adjustment of the feedback or tap on the coil is inconvenient. For the Colpitts, the tank capacitor shaft is at ground potential which prevents r-f burns to the operator. Also, the feedback voltage ratio can be adjusted easily by employing separate variable capacitors. Fixed capacitors are used in some Colpitts tank circuits and the inductance of the tank coil is varied. Other models use plug in coils to cover a wide range of frequencies without disturbing the feedback voltage ratio established by the fixed tank capacitors.

As the operating frequency in a vacuum tube circuit increases, the reactance of the tube interelectrode capacitance becomes important. This capacitance plus other effects in the tube limit the maximum operating frequency. The interelectrode capacitance can be put to use in a high-frequency oscillator similar to the Colpitts oscillator. Figure 13B shows such an oscillator, called an **ULTRAUDION OSCILLATOR**. The plate-cathode capacitance (C_{pk}) and grid-cathode capacitance (C_{gk}) of the tube form the tapped tank capacitance. C_1 is included in Figure 13B to permit the tank circuit to be tuned to the desired frequency. C_3 and R_1 form the shunt grid leak bias network. C_2 couples the plate of V_1 to the tank circuit.

Aside from the fact that the tapped tank circuit capacitance is the tube interelectrode capacitance, the ultraudion oscillator operates exactly like the Colpitts oscillator. As mentioned earlier, it is used in high-frequency applications where the tube capacitance is high enough to provide the desired feedback. This type of oscillator is often used in the tuner circuits of TV receivers.

THE TUNED-PLATE TUNED-GRID OSCILLATOR

Another basic type of oscillator circuit, shown in Figure 14, includes tuned circuit L_1C_1 connected to the grid, and tuned circuit L_2C_3 connected to the plate. Therefore, the arrangement is known as a **TUNED-PLATE TUNED-GRID OSCILLATOR (TPTG)**.

Compared with the circuit of Figure 10, there is no apparent coupling between the plate and grid coils. However, to maintain the conditions required for oscillation, there must be some form of feedback. Actually, in the tuned-plate tuned-grid oscillator the circuits are coupled by the grid-plate capacitance of the tube, as indicated by the broken line capacitor C_{gp} in Figure 14.

To obtain proper operation, the plate tuned circuit, C_3L_2 , is tuned slightly higher than the grid circuit, L_1C_1 . The amount of feedback can be varied

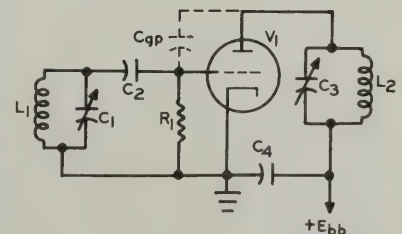


Figure 14

by adjusting the tuning of either the grid or plate circuits. The frequency of oscillation is determined by the tuned circuit that has the highest Q.

THE ELECTRON-COUPLED OSCILLATOR

In radio transmitters and receivers, and in many other applications, the oscillator frequency must be held extremely constant. Undesired changes of oscillator frequency result from three major causes: (1) changes in the mechanical arrangement of the elements in the oscillating circuit; (2) variations in the amplification factor, grid and plate resistances, and interelectrode capacitances of the tube; and (3) changes in various circuit element values.

Changes in the mechanical arrangement of the circuit components may be produced by vibration, temperature variations, or by electrostatic and electromagnetic forces. When necessary, these can be minimized by careful electric and mechanical design and by temperature control.

Variations in the amplification factor, grid and plate resistances, and interelectrode capacitances of a tube are caused by changes in the voltages applied to the electrodes and by changes in cathode emission. To prevent these changes, the operating voltages may be stabilized by voltage regulating devices. The effect of changing interelectrode capacitance on oscillator frequency may be minimized by employing a high ratio of tuning capacitance to inductance (high C to L ratio) and by using circuits in which the tuning capacitance shunts the tube grid-plate capacitance. With this arrangement, the tube interelectrode capacitance is but a small part of the total tuning capacitance, and any variation in it will cause only a slight frequency change.

Changes in the circuit elements are caused by changes in the temperatures of inductors and capacitors, and by variations of load, which alter the effective ac resistance of the tuned circuits. Changes of inductance and capacitance may be minimized by the careful choice of apparatus and the proper location of component parts, temperature control, and the use of temperature controlled compensating capacitors.

The oscillation frequency of an LC circuit depends not only on its inductance and capacitance, but also on the effective ac resistance of the inductor, which in turn depends on the power delivered by the inductor to the load.

The frequency of oscillation is therefore affected by load changes unless the output is taken from the oscillator in such a manner that the current in the inductor is not changed. It was for this purpose that the ELECTRON-

COUPLED OSCILLATOR (ECO) shown in Figure 15 was developed. Here, the only coupling between the oscillator and the load is through the electron stream of the tube. Due to the "electron coupling", changes in the load have very little effect on the frequency controlling circuits, and so the oscillator frequency remains relatively constant.

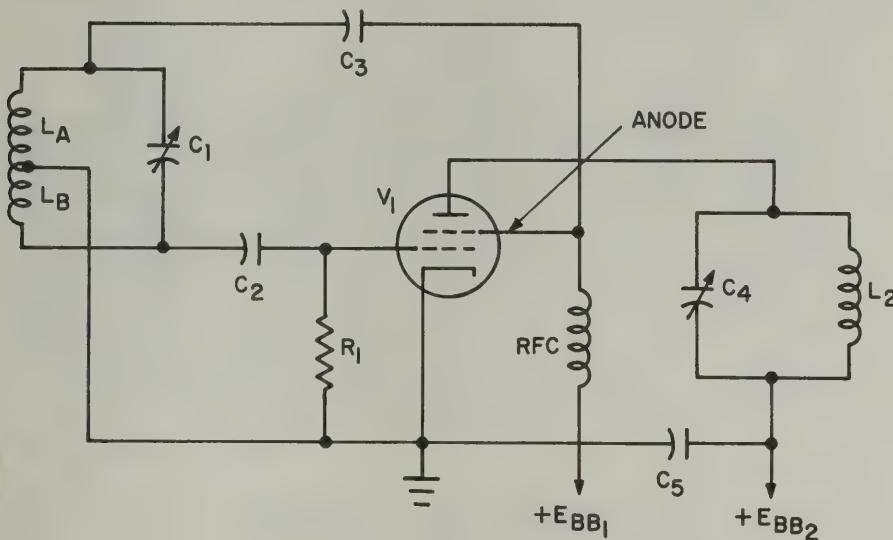


Figure 15

The electron-coupled oscillator frequency is determined by the inductance and capacitance in tank circuit $L_A L_B C_1$. The circuit differs from those explained previously in that no capacitive, inductive, or direct coupling is used between the tank circuit and the desired load.

Instead, it employs a tetrode or pentode tube in which the screen grid serves as the plate of a triode self-excited oscillator. To prevent confusion, the screen grid in this circuit is labeled "anode". Thus, the cathode, control grid and anode may be connected as a triode tube in a regular oscillator circuit such as an Armstrong, Hartley, or Colpitts.

Neglecting for a moment the plate circuit in Figure 15, the remaining tube elements, along with circuit components L_A , L_B , C_1 , C_2 , C_3 , R_1 , and RFC constitute a shunt-fed Hartley oscillator circuit just like Figure 12. However, in this arrangement the screen grid of tube V_1 functions as the oscillator anode. For this part of the circuit, any other oscillator circuit could be used just as well as the Hartley.

Since the screen grid or anode of the tube is made of a mesh material, it captures only part of the emitted electrons, while the rest continue on to the plate. Thus, only enough electrons to supply the driving power for the oscillator need be taken by the anode.

The remaining electrons pass through the anode in pulses at the frequency of oscillation in the $L_A L_B C_1$ tank circuit, and this pulsating current supplies the power for the plate-circuit load $C_4 L_2$.

An important advantage of the ECO is that an increase in plate voltage shifts the frequency in one direction, while an increase in anode voltage shifts it in the opposite direction. Thus, by supplying the anode through a voltage divider arrangement in the plate supply, the anode voltage can be adjusted until the frequency is independent of the plate supply voltage. Under these conditions, any change in frequency due to a slight change in plate voltage is nullified by an equal and opposite change in frequency due to the change in anode voltage. In general, the stability is best when the ratio of the plate to anode voltage is about three to one.

THE CLAPP OSCILLATOR

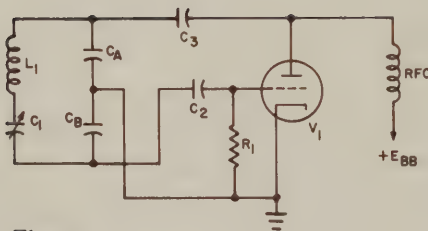
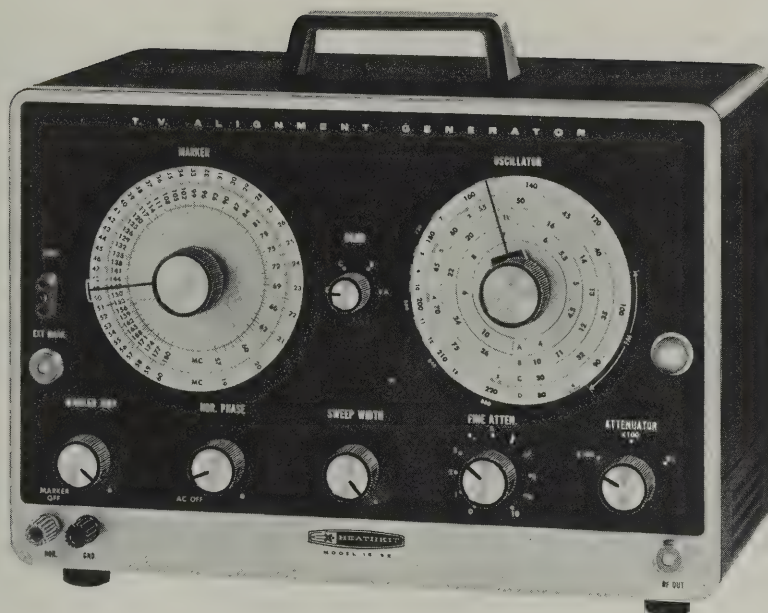


Figure 16

Figure 16 shows a circuit called a **CLAPP OSCILLATOR**. Except for the added capacitor C_1 in series with tank coil L_1 , this circuit is like the Colpitts oscillator of Figure 13. As in the Colpitts oscillator, C_A and C_B divide the tank voltage. The voltage across C_B is fed to the grid of V_1 through C_2 as the feedback voltage.

The dc plate voltage of V_1 is blocked from the tank circuit by C_3 , hence the Clapp oscillator is shunt-fed. Class C operation is obtained through the $R_1 C_2$ grid leak bias network.

Basically, the oscillating action is the same as in the Colpitts oscillator. However, the frequency stability of the Clapp oscillator is better than in a Colpitts oscillator. This stability is obtained by the use of series capacitance C_1 . The capacitance of C_1 is less than the equivalent series capacitance of C_A and C_B . As a result, the oscillation frequency depends mainly on L_1 and C_1 . At the resonant frequency of the tank, the net reactance of L_1 and C_1 in series is low. As a result, the oscillation frequency is not affected to any extent by changes in the characteristics of V_1 . Adjustment of C_1 selects the desired frequency.



This TV alignment generator includes a crystal controlled oscillator, a variable frequency LC oscillator and an oscillator that can be swept over a band of frequencies.

Courtesy Heath Co.

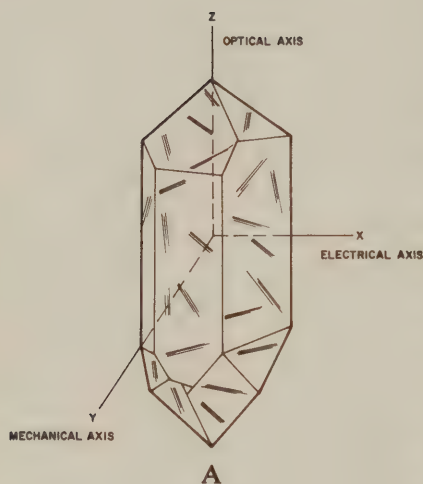
THE CRYSTAL OSCILLATOR

Another method of achieving a high degree of frequency stability is to employ the electromechanical oscillation properties of a quartz crystal. When a properly cut plate of a quartz crystal is placed under mechanical strain, electric charges appear on opposite faces. Also, when electric charges are applied to opposite faces, mechanical deformation of the plate occurs. This action is known as the **PIEZOELECTRIC EFFECT**.

An alternating voltage applied across the crystal causes it to vibrate. If the frequency of the applied voltage approximates a frequency at which mechanical resonance can occur in the crystal, the amplitude of the vibrations is large. Advantage is taken of this effect by mounting the crystal plate in a suitable holder and employing it in place of the tuned circuit that normally controls the frequency of the oscillator. When properly adjusted, the circuit oscillates at the mechanical resonant frequency of the crystal plate and is affected very little by circuit changes.

CRYSTAL CUTS

The crystal plates employed in oscillator circuits are cut from whole or "mother" crystals which have the general appearance shown in the sketch



of Figure 17. As indicated by the broken lines, the crystal has three major axes, with the letter "X" used to denote what is known as the electric axis, "Y" the mechanical axis, and "Z" the optical axis.

For use in oscillator circuits, small pieces or plates are cut from the mother crystal and mounted in a holder like that shown in Figure 17B. The resonant or natural vibration frequency of these plates depends on their dimensions, the type of mechanical oscillation involved, and the orientation of the plates as they are cut from the mother crystal. Usually, the frequency is inversely proportional to the thickness; that is, the thinner the plate, the higher the frequency, and the thicker the plate, the lower the frequency.

There are several cuts which can be made from a crystal. A plate cut with its larger surfaces perpendicular to the X axis is known as an "X-cut" plate. A "Y-cut" plate is cut so that its major surfaces are perpendicular to the Y axis.

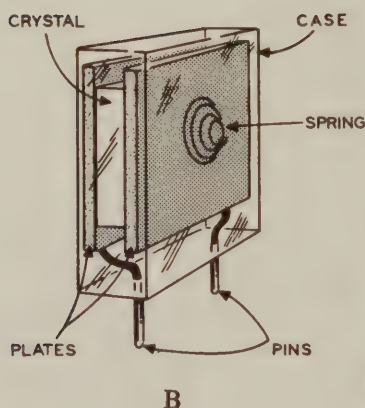


Figure 17

TEMPERATURE COEFFICIENT OF CRYSTALS

Although the oscillation frequency depends primarily on the physical dimensions of the crystal plate, the type of cut and the mode of mechanical vibration, the natural resonant frequency varies slightly due to slight internal changes caused by heat. The **TEMPERATURE COEFFICIENT (TC)** of a quartz crystal is the fractional change in frequency per unit change of temperature and may be negative, positive, or zero.

THE TEMPERATURE COEFFICIENT IS NEGATIVE, as in an X-cut crystal plate, WHEN THE FREQUENCY CHANGE VARIES INVERSELY AS THE TEMPERATURE CHANGE. If the surrounding temperature increases, the frequency of an X-cut crystal decreases, and if the temperature decreases, the frequency increases.

THE TEMPERATURE COEFFICIENT IS POSITIVE, as in a Y-cut crystal, WHEN THE FREQUENCY CHANGE VARIES DIRECTLY AS THE TEMPERATURE CHANGE.

Other cuts, known as the GT and the doughnut types, have zero temperature coefficients within certain temperature ranges. Changes of temperature within these ranges have no effect on the frequency.

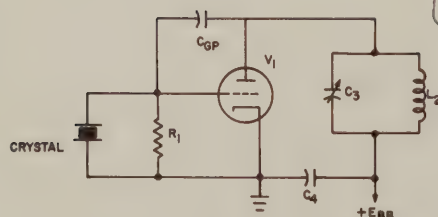


Figure 18

CRYSTAL OSCILLATOR CIRCUITS

Figure 18 shows the circuit diagram of a **CRYSTAL CONTROLLED OSCILLATOR** in which the frequency controlling element is a piezoelectric

crystal. This circuit is similar to that of the tuned-plate tuned-grid oscillator of Figure 14 except that the crystal replaces the L_1C_1 grid circuit. In both cases, feedback is obtained through the grid-plate capacitance C_{gp} of the tube. As suggested by the diagram symbol, the crystal plate is mounted between metal plates to which the circuit connections are made. The crystal holder plates also function as the plates of a capacitor connected from ground to the V_1 grid. During operation, this capacitor charges through the tube when the V_1 grid is positive with respect to the cathode, and discharges through resistor R_1 when the grid is negative. Thus, the crystal holder capacitance acts just like that of C_2 in Figure 14 and maintains a constant negative bias on the V_1 grid.

When the tube is placed in operation, the initial surge of plate current produces a voltage pulse. This pulse, due to the interelectrode capacitance C_{gp} , is impressed across the crystal and causes it to oscillate at its natural or resonant frequency. As the crystal vibrates, it generates a voltage of like frequency which is impressed across the grid circuit. Like the other oscillators, the amplifying action of the tube causes the changes of grid voltage to reappear with greater magnitude in the plate circuit. The changes of plate voltage are fed back through the grid-plate capacitance with sufficient amplitude to maintain the vibration of the crystal, and thus the oscillations continue at a constant amplitude.

As in the tuned-plate tuned-grid oscillator, the plate circuit must be tuned to a frequency that is slightly higher than that of the crystal to provide feedback in the proper phase. However, the frequency of oscillation is controlled by the vibration of the crystal, and within ordinary limits it is independent of the other electric circuit conditions.

Figure 19 shows the circuit diagram of a type of oscillator that combines the advantages of crystal control and electron coupling. This circuit is somewhat similar to the electron-coupled oscillator of Figure 15. The feedback path is from the cathode through the crystal circuit to the control grid.

The circuit is called a **TRI-TET OSCILLATOR**, because the cathode, control grid, and screen grid of a tetrode or four-element tube are used as a triode oscillator circuit. This oscillator includes all the good features found in Hartley and electron coupled circuits along with crystal control. Consequently, it is a very popular circuit.

To prevent unstable operation, it is common practice to use a separate power supply for the crystal oscillator. If a common power supply with poor regulation is used, any voltage fluctuations caused by changing load conditions

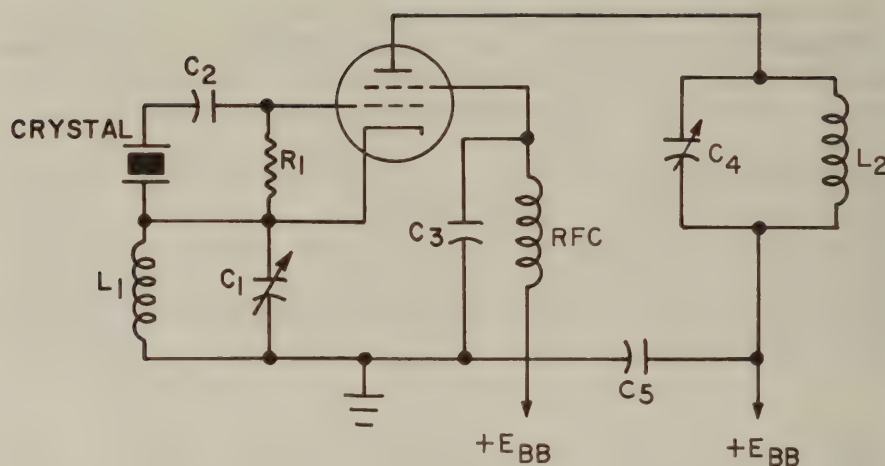


Figure 19

result in changing plate and screen voltages. This varies the frequency characteristics of the crystal.

The excellent performance characteristics of the compact crystal oscillators fully justify the additional precautions that must be taken with them. The equivalent tank circuit formed by the crystal has an exceptionally high Q which results in excellent frequency stability, especially when a constant temperature is maintained. Q 's in excess of 1,000, which are considerably greater than those obtained with good LC tank oscillator circuits, are commonplace.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

THE COLPITTS OSCILLATOR

THE TUNED-PLATE TUNED-GRID OSCILLATOR

THE ELECTRON-COUPLED OSCILLATOR

THE CLAPP OSCILLATOR

THE CRYSTAL OSCILLATOR

CRYSTAL CUTS

TEMPERATURE COEFFICIENT OF CRYSTALS

CRYSTAL OSCILLATOR CIRCUITS

27. How does a Colpitts oscillator differ from a Hartley?
28. What happens to the amount of feedback voltage in the circuit of Figure 13 if the value of C_B is increased?
29. How is feedback obtained in a TPTG oscillator?
30. In the TPTG oscillator circuit, both the grid and plate circuits are tuned to exactly the same frequency. True or False?
31. Name some of the factors which affect the stability of an LC oscillator.
32. How is the load coupled to the oscillator circuit in an ECO?
33. Can frequency changes in an ECO due to plate voltage variations be eliminated?
34. How does a Clapp oscillator differ from a Colpitts?
35. In the circuit of Figure 16, what components determine the frequency of oscillation?
36. What is meant by the term "temperature coefficient" with respect to crystals?
37. How does a crystal with a negative temperature coefficient differ from one with a positive temperature coefficient?
38. What type of oscillator circuit is similar to the circuit of Figure 18?

39. In Figure 18, is the L_2C_3 tank circuit tuned to a higher or lower frequency than that of the crystal?
40. Is the circuit of Figure 19 similar to that of the Colpitts circuit?
41. Why is it good practice to operate an oscillator circuit from a separate regulated plate supply?

PARASITIC OSCILLATIONS

The frequency at which an oscillator operates is normally controlled by the inductance and capacitance of the tank circuit, or by the natural resonant frequency of a crystal. However, the interelectrode capacitances of the tube circuit may produce resonance in conjunction with the grid and plate circuit leads. These wires can be thought of as one-turn inductors which produce oscillations at frequencies much higher than the resonant frequency of the tank circuit or crystal.

This action produces what are known as **PARASITIC OSCILLATIONS**. These oscillations are generally at a frequency different from the operating frequency and are very undesirable because they absorb power that, in turn, robs the available power output of the oscillator or h-f amplifier involved.

Referring to Figure 14 as an example, at very high frequencies capacitors C_3 and C_4 offer negligible reactance. This may permit the inductance of the connecting wires from the plate to C_3 , from C_3 to C_4 , and from C_4 to the cathode to form a tank circuit with the plate-cathode capacitance of the tube. In like manner, the grid circuit may resonate at a frequency determined by the grid-cathode capacitance and the inductance of the connecting wires from the grid to C_2 , C_2 to C_1 , and from C_1 to the cathode. Thus, the complete circuit can act as a tuned-plate tuned-grid oscillator which resonates at some very high frequency.

As mentioned earlier in this lesson, to obtain feedback voltage of the proper phase, the plate tank of a tuned-plate tuned-grid oscillator must be tuned to a frequency slightly higher than that of the grid circuit. To prevent parasitic oscillations which are a result of TPTG oscillator action, the plate circuit inductance is increased until the circuit resonates at a frequency lower than that of the grid circuit.

Often this is done by increasing the length of the plate circuit connecting wires or leads. In many cases the extra length is wound into a coil of several turns which then is called a parasitic choke. At the proper operating frequency, the reactance of this choke is too small to cause any appreciable effect.

Another system of suppressing parasitic oscillations is to insert small resistors in the tube plate or grid leads. Since these resistors are inserted directly into the parasitic frequency tuned circuit, they lower the tank circuit Q until oscillations cannot be sustained. On the other hand, since the resistors are not in the operating frequency tank circuit, but rather in the lead going to it, there is no appreciable effect on the desired operation of the oscillator.

As a specific example, such a resistor would be inserted between the plate and C_3 or between the grid and C_2 of Figure 14. In the latter case, for instance, although not in the operating frequency tank circuit L_1C_1 , the resistor would be directly in series with the parasitic oscillation tank circuit as previously traced. Hence, the lowered Q of this circuit would not sustain the parasitic oscillations.

RC OSCILLATORS

RC oscillators employ resistors and capacitors in their feedback networks. As in LC oscillators, the feedback network in an RC oscillator must provide a feedback signal of the proper phase and amplitude to sustain oscillations. In most cases, RC oscillators are used to develop signals in the low frequency range, below about 100 Hz. LC oscillators, on the other hand, are used to develop higher frequency signals.

THE MULTIVIBRATOR

One of the most common RC oscillators is the **MULTIVIBRATOR**. It is nothing more than a two-stage resistance-capacitance coupled amplifier with the output fed back to the input by C_2 as shown in Figure 20. Such an arrangement oscillates because each tube produces a phase shift of 180 degrees.

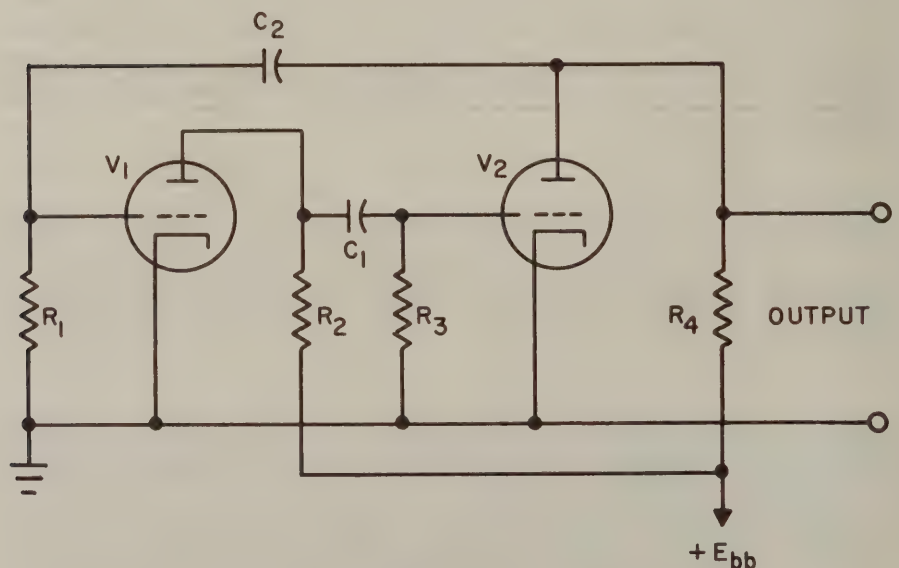


Figure 20

This causes the second tube to supply the first tube with an input voltage that has the right phase to sustain oscillation. A multivibrator develops a nonsinusoidal output which is rich in harmonics. The frequency of a multivibrator may be readily controlled by a voltage applied to one of the grids.

Because of its numerous applications, the multivibrator circuit has many variations. However, the general principles of operation are the same for all types. The following explanations of the actions in the basic circuit will aid in the understanding of any modification which may be encountered.

The multivibrator operates in a see-saw manner with one tube conducting while the other is cut off. The circuit remains in a stable state for a definite period of time and then changes so the conducting tube cuts off and the cutoff tube conducts. This reversal action continues as long as power is applied. To examine circuit operation let's assume that V_1 has just gone into cutoff and V_2 into conduction. Since V_2 was originally cut off, C_2 was charged to approximately $+E_{bb}$. C_2 now discharges through R_1 driving the grid of V_1 negative, holding it cut off. With V_1 cut off, C_1 is charging to $+E_{bb}$ through R_3 and the grid cathode circuit of V_2 driving the grid of V_2 positive, holding it in conduction.

This temporarily stable condition exists until capacitor C_2 discharges enough through resistor R_1 to permit the V_1 grid voltage to decrease sufficiently to allow a small plate current in tube V_1 . The V_1 plate current produces a small voltage drop across plate load resistor R_2 , thereby lowering the V_1 plate voltage.

With reduced V_1 plate voltage, capacitor C_1 starts to discharge through resistor R_3 , and the resulting voltage drop across R_3 drives the grid of tube V_2 negative with respect to its cathode. With a negative grid tube V_2 conducts less, and the resultant rise of V_2 plate voltage causes capacitor C_2 to charge through resistor R_1 and the V_1 grid cathode circuit and impress a positive voltage on the grid of tube V_1 .

The positive grid greatly increases the V_1 plate current, thus causing the plate voltage to drop and permit a further discharge of capacitor C_1 . Carried by resistor R_3 , this discharge current develops a voltage drop that drives the V_2 grid sufficiently negative to hold the tube cut off for a period of time while capacitor C_1 discharges.

By the time the C_1 discharge is nearly complete, the drop across the grid resistor decreases to less than the cutoff voltage, and tube V_2 becomes conductive. The V_2 plate current produces a voltage drop across plate load re-

sistor R_4 , thus causing the V_2 plate voltage to decrease. This condition forces capacitor C_2 to discharge through grid resistor R_1 and makes the V_1 grid negative with respect to its cathode.

With a negative grid, tube V_1 conducts less, and the resulting rise of plate voltage causes capacitor C_1 to charge through grid resistor R_3 and the V_2 grid cathode circuit and impress a positive voltage on the grid of tube V_2 . The positive grid permits the V_2 plate current to rise to a high value, and the corresponding decrease of plate voltage permits capacitor C_2 to discharge through resistor R_1 sufficiently to bias tube V_1 to cutoff.

This condition continues until capacitor C_2 is almost discharged and the drop across resistor R_1 is slightly less than the cutoff voltage. At this point, tube V_1 becomes conductive to begin the next cycle. Thus, the entire action is repeated over and over, with tube V_1 cut off when tube V_2 is conducting, and tube V_1 conducting when tube V_2 is cut off.

The portion of each cycle that tube V_1 is cut off depends upon the R_1C_2 time constant, and the V_2 cutoff interval depends upon the R_3C_1 time constant. In the circuit of Figure 20, the grid voltages have the waveforms pictured in Figure 21A and the plate voltages have the waveforms pictured in Figure 21B when $V_1 = V_2$, $R_1 = R_3$, $R_2 = R_4$, and $C_1 = C_2$. However, the V_1 grid and plate voltages are 180 degrees out-of-phase with the V_2 voltages. Either grid or plate voltage may be taken as the output of the circuit.

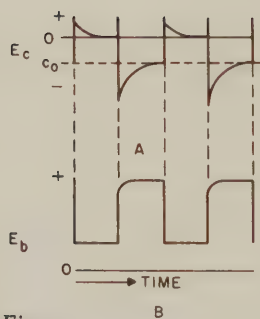


Figure
21

Because of the symmetry of the components and connections, the Figure 20 circuit plate output has positive and negative pulses of equal duration; therefore, it is called a **SYMMETRICAL MULTIVIBRATOR**.

An **ASYMMETRICAL MULTIVIBRATOR** is one in which unequal components are employed to produce narrow and wide pulses. To illustrate this arrangement, assume that $V_1 = V_2$, $R_1 = R_3$, $R_2 = R_4$, and $C_1 = 3C_2$. Under this condition, the R_3C_1 time constant is three times as great as the R_1C_2 time constant. Consequently, tube V_2 is cut off approximately three times as long as tube V_1 . As a result, the V_2 output voltage positive pulses have a time duration or pulse width about three times as great as its negative pulses.

In a like manner, an output from the plate of V_2 with positive pulses whose time duration is about one-third as long as the negative pulses may be obtained by making the C_2 capacitance three times as great as that of C_1 while leaving the other parts unchanged.

THE CATHODE-COUPLED MULTIVIBRATOR

Another form of multivibrator circuit is shown in Figure 22. Although separate tubes are shown, a duo-triode may be used. As shown, the V_1 plate is coupled to the V_2 grid through capacitor C_1 , but there is no coupling from the V_2 plate back to the V_1 grid. Instead, the necessary feedback is obtained by means of the common unbypassed cathode resistor, R_1 , which carries both plate currents. Any variations of either plate current produce corresponding variations of bias voltage E_{R_1} .

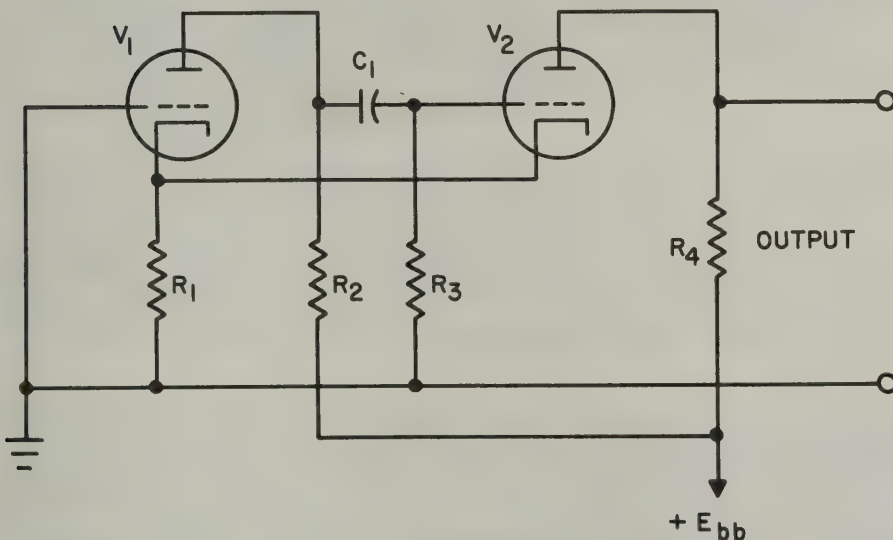


Figure 22

An important point to observe is that the V_1 grid to cathode circuit contains only resistor R_1 , while the V_2 grid to cathode circuit contains resistors R_1 and R_3 . Thus, the V_1 grid bias is equal to E_{R_1} , and the V_2 grid bias is equal to the sum of or the difference between E_{R_1} and E_{R_3} , depending on their relative polarities at any instant. Another point to remember is that in practice resistor R_3 has a much higher resistance than R_1 .

Oscillation is started by a very small change of voltage on the grid of one of the tubes, for example, a positive change on the V_2 grid. The resulting increase of V_2 plate current causes a greater voltage drop across cathode resistor R_1 that in turn tends to lower both plate currents. The decrease in V_1 plate current results in less voltage drop across resistor R_2 . Thus, the V_1 plate voltage rises, and capacitor C_1 charges through resistor R_3 and produces voltage E_{R_3} which is positive at the grid end of R_3 . At this instant, E_{R_1} and

E_{R_3} oppose each other so that the net bias on V_2 decreases even though E_{R_1} increases.

With less bias, tube V_2 conducts more heavily, a greater negative bias is applied to tube V_1 , and tube V_2 is driven into heavier conduction. This action continues rapidly until E_{R_3} is greater than E_{R_1} , at which time the V_2 grid is driven positive with respect to the cathode and draws current. Carried by resistor R_1 , the V_2 plate and grid currents develop a voltage which is high enough to bias V_1 beyond cutoff.

This temporarily stable condition exists until the C_1 charge current decreases enough to permit the V_2 grid to become negative with respect to the cathode and stop the grid current. By itself, the V_2 plate current cannot develop enough bias across resistor R_1 to keep tube V_1 cut off, so the tube begins to conduct.

The resulting voltage drop across resistor R_2 lowers the V_1 plate voltage and permits capacitor C_1 to discharge through resistor R_3 . This time, E_{R_1} and E_{R_3} aid each other and increase the net negative bias on tube V_2 . With a higher negative bias, tube V_2 conducts less, a lower voltage is developed across resistor R_1 , and tube V_1 conducts more heavily. This action continues rapidly until the V_2 grid is driven negative enough to cut off the tube.

The circuit remains in this temporarily stable state until capacitor C_1 has discharged enough through resistor R_3 to permit the grid to rise above the cutoff voltage and initiate plate current in tube V_2 , thus starting a new cycle. As in the circuit of Figure 20, the output may be taken from either plate or from the grid of the second tube. The plate output connection is pictured in Figure 22.

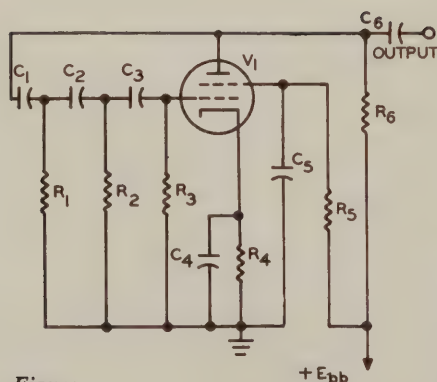


Figure 23

THE PHASE SHIFT OSCILLATOR

One of the simplest RC oscillators is the **PHASE SHIFT OSCILLATOR** shown in Figure 23. V_1 operates as a conventional amplifier. R_4 and C_4 provide cathode bias. R_5 is a screen dropping resistor and C_5 is a screen bypass capacitor. R_6 is the plate load resistor and C_6 couples the output signal from the V_1 plate. Feedback from the plate of V_1 back to its grid is provided by C_1R_1 , C_2R_2 and C_3R_3 .

A common cathode amplifier provides a 180° phase shift between its input and output signals. The feedback network shifts the phase of the feedback signal so it is in phase with the input signal, thus causing the amplifier to oscillate.



Audio generators like the one shown above employ RC oscillators to develop a sinusoidal output signal. This unit has a range from 1 Hz to 100 kHz.

Courtesy Eico Electronic Instrument Co., Inc.

In the circuit of Figure 23, the voltage at the plate of V_1 is 180° out-of-phase with the signal at its grid. You may recall that the voltage across a resistor in a series RC network is out of phase with the supply voltage. In Figure 23, R_1 and C_1 are in series with the plate of V_1 . The values of R_1 and C_1 are chosen so the voltage across R_1 is 60° out-of-phase with the V_1 plate voltage at the oscillator frequency. This voltage is applied to the series combination of C_2 and R_2 . The values of R_2 and C_2 are chosen so the voltage across R_2 is 60° out of phase with the voltage across R_1 at the oscillator frequency. The values of C_3 and R_3 are also chosen to provide a 60° phase shift.

Since each series RC circuit is designed to provide a 60° phase shift at the oscillator frequency, the voltage applied to the grid of V_1 is the plate voltage of V_1 shifted $60^\circ + 60^\circ + 60^\circ$, or 180° . The V_1 plate voltage is 180° out-of-phase with its grid voltage. With the 180° phase shift provided by the RC feedback network, the feedback signal applied to the V_1 grid is $180^\circ + 180^\circ = 360^\circ$ or in phase with the grid voltage. Thus, oscillations occur.

The phase shift oscillator shown in Figure 23 produces a sinusoidal output signal. The frequency of oscillation is determined by the values in the C_1R_1 , C_2R_2 and C_3R_3 networks. Since each RC network provides the same phase

shift, all the capacitors have the same value and all the resistors have the same value.

The LC oscillators described earlier used grid leak bias to provide Class C operation. A normal Class C amplifier produces the most distortion even though it has the highest efficiency. Due to the characteristics of the tank circuit, an LC oscillator develops a sinusoidal output waveshape with low distortion even though the tube operates Class C.

Class A operation is used in the phase shift oscillator of Figure 23 to obtain a sinusoidal output waveshape with low distortion. This is because there is no LC tank circuit to correct distortion in the phase shift oscillator.

The frequency of the phase shift oscillator of Figure 23 can be changed by changing the values of all the resistors or capacitors in the feedback circuit. The phase shift oscillator is often used for fixed frequency applications.

THE WEIN BRIDGE OSCILLATOR

Another widely used RC oscillator is the **WEIN BRIDGE OSCILLATOR**. This oscillator makes use of a special RC circuit called the Wein bridge. Figure 24 shows the circuit of a Wein bridge oscillator. V_1 and V_2 operate as conventional amplifiers. Both V_1 and V_2 provide a 180° phase shift be-

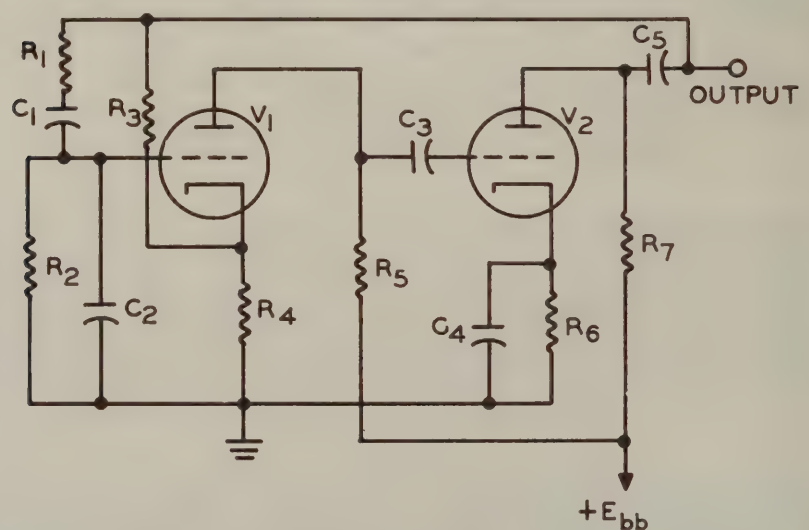


Figure 24

tween their grid (input) and plate (output) signals. As a result, the signal at the output terminal, the plate of V_2 , is in phase with the signal at the grid of V_1 . According to the principles illustrated in Figures 1 and 2, if the signal at the plate of V_2 were coupled to the grid of V_1 , the circuit would oscillate.

In Figure 24, the signal at the plate of V_2 is coupled back to the grid of V_1 through the RC network consisting of R_1C_1 and R_2C_2 . R_1C_1 and R_2C_2 together with R_3 and R_4 form the Wein bridge. The voltage appearing across R_2 and C_2 is in phase with the signal coupled to the network from the plate of V_2 at only one frequency. This is the frequency at which the circuit oscillates.

Like the phase shift oscillator, the Wein bridge oscillator develops a sinusoidal output wave. The frequency of oscillation is controlled by the values of R_1C_1 and R_2C_2 . For proper circuit operation the values are chosen so the product of R_1 and C_1 equals the product of R_2 and C_2 ; that is $R_1C_1 = R_2C_2$. For this reason, both resistors or both capacitors must be changed at the same time whenever the oscillation frequency is to be changed. Ganged resistors and capacitors are commonly used in this circuit to change the frequency.

The Wein bridge circuit is widely used as a variable frequency test oscillator. One example is its use as an "audio generator" to supply test signals for testing hi-fi stereo music systems.

THE BLOCKING OSCILLATOR

Figure 25 shows the circuit of an RC oscillator called a **BLOCKING OSCILLATOR**. Although this circuit resembles the LC oscillator shown in Figure 6, it does not employ a tank circuit nor does it develop a sinusoidal output signal. Feedback from the plate of V_1 to its grid is provided by T_1 as in an LC oscillator. C_1 and R_1 form an RC network which controls the frequency of oscillation.

Let us examine the circuit operation when power is first applied to the circuit. With no bias on V_1 , its plate current rapidly increases to maximum as power is applied. The increasing V_1 plate current in the primary of T_1 induces a feedback voltage in the T_1 secondary. The secondary leads are connected so that the grid of V_1 is driven positive at this time. C_1 is charged by V_1 grid current to the polarity shown.

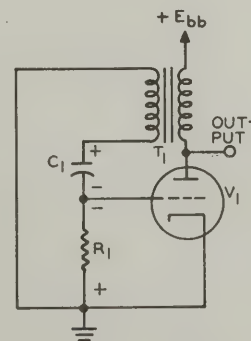


Figure
25

When V_1 plate current reaches its maximum saturated value, a voltage is no longer induced into the T_1 secondary because the primary current is not

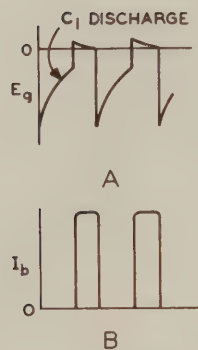


Figure
26

changing. With no voltage at the secondary of T_1 , C_1 begins to discharge through R_1 and the T_1 secondary. C_1 discharge current develops a voltage across R_1 with the polarity shown. This voltage drives V_1 into cutoff. V_1 remains cut off until the voltage across R_1 is less than that required to cut the tube off. Then V_1 conducts again and a new cycle of oscillation begins.

Since V_1 conducts for only short intervals, its plate current consists of pulses as shown in Figure 26B. The V_1 grid voltage has the waveform shown in Figure 26A. The blocking oscillator is sometimes called a "pulse generator" because it develops an output signal consisting of a series of pulses. As shown, the output signal is usually taken from the plate of V_1 .

The frequency of oscillation in the blocking oscillator of Figure 25 is controlled by the values of R_1 and C_1 . Increasing the value of either R_1 or C_1 decreases the frequency, and vice versa. The grid voltage waveform shown in Figure 26A is characteristic of the blocking oscillator. This waveform is sometimes referred to as a "cash register" waveform.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

PARASITIC OSCILLATIONS

RC OSCILLATORS

THE MULTIVIBRATOR

THE CATHODE-COUPLED MULTIVIBRATOR

THE PHASE SHIFT OSCILLATOR

THE WEIN BRIDGE OSCILLATOR

THE BLOCKING OSCILLATOR

42. What are parasitic oscillations?
43. How do low value resistors in the plate and grid leads of a tube prevent parasitic oscillation?
44. Does a multivibrator develop a sinusoidal output signal?
45. What components in the circuit of Figure 20 determine how long V_1 remains cut off?
46. If the values of C_1 and C_2 in Figure 20 are decreased, what happens to the frequency?
47. How does the output of a symmetrical oscillator differ from that of an asymmetrical oscillator?
48. How is feedback from the V_2 stage to the V_1 stage in the circuit of Figure 22 obtained?
49. Can an output signal be taken from the grid of V_1 in Figure 22?
50. Does C_1 in Figure 22 charge or discharge when V_1 is conducting?
51. Does the phase shift oscillator circuit develop a sinusoidal output signal?
52. How is the positive feedback necessary to sustain oscillation developed in the circuit of Figure 23?
53. Would the circuit of Figure 23 be used to generate radio frequencies?
54. In the circuit of Figure 24, what is the relation between R_1C_1 and R_2C_2 ?
55. Over what frequency range does the circuit of Figure 24 operate?

56. The blocking oscillator develops a sinusoidal output signal. True or False?
57. What components in the circuit of Figure 25 control the frequency of oscillation?
58. What type of signal is available at the grid of V_1 in Figure 25 ?

SUMMARY

An electronic oscillator is a device that generates an alternating current. It is basically an amplifier with its output fed back to the input with the proper phase and amplitude to sustain oscillations.

Oscillators may be classed according to their output waveshape. LC oscillators usually have a sinusoidal output while RC oscillators, such as the multivibrator, have a nonsinusoidal output.

The LC oscillator contains a tank circuit which determines the frequency of oscillation. The electron tube in this type of oscillator acts as a switch to replace the energy lost by the tank circuit. Class C operation in conjunction with grid leak bias is employed to secure maximum efficiency.

The Hartley and Colpitts oscillators provide feedback with tapped tank circuit components, while the Armstrong oscillator uses a tickler coil to obtain feedback. The Clapp oscillator is a type of Colpitts oscillator employing series tuning of the tank inductor. The TPTG oscillator employs tuned tank circuits in the grid and plate leads with the grid-to-plate capacitance of the tube providing feedback. If the dc plate current flows through part of the tank coil, the oscillator is series-fed. In shunt-fed oscillators the dc plate current does not flow through the tank circuit.

To prevent loading and frequency changes, electron coupling may be employed. This arrangement couples the output of the oscillator to the load through the tube's electron stream. When high frequency accuracy and stability are required, crystal controlled oscillators are employed. These oscillators make use of the piezoelectric effect of quartz crystals.

Nonsinusoidal oscillators such as the multivibrator are usually RC coupled feedback amplifiers. In the plate coupled multivibrator, feedback from the output of the second tube through an RC network to the grid of the first tube sustains oscillation. The cathode coupled multivibrator develops feedback through the common cathode resistor. The pulse duration and operating frequency of a multivibrator are determined by the RC networks in the grid circuits of the tubes.

The phase shift of Wein bridge RC oscillators develops a sinusoidal output signal. These circuits are generally used in the low (audio) frequency range. The blocking oscillator is another RC oscillator that develops a nonsinusoidal output. The circuit of the blocking oscillator is very similar to that of the Armstrong circuit except that it is controlled by an RC circuit.

IMPORTANT DEFINITIONS

ARMSTRONG OSCILLATOR — A form of feedback oscillator in which the tuned tank is in the grid circuit of the tube, and the feedback coupling is in the plate circuit.

BLOCKING OSCILLATOR — A nonsinusoidal RC oscillator employing a transformer for feedback and an RC circuit to determine its operating frequency.

CLAPP OSCILLATOR — A type of Colpitts oscillator in which the inductor is tuned by a series capacitor. The two capacitors across the tank serve as the feedback network.

COLPITTS OSCILLATOR — A form of feedback oscillator in which two capacitors are connected in series across the tank inductor, from plate to grid and with the junction between them connected to the cathode.

CRYSTAL CONTROLLED OSCILLATOR — A type of feedback oscillator in which the frequency controlling element is a piezoelectric crystal.

ELECTRON-COUPLED OSCILLATOR — (ECO) — A type of oscillator in which the electrode connected to the load is coupled to the oscillator proper only by the electron stream within the tube.

HARTLEY OSCILLATOR — A form of feedback oscillator in which a tapped tank circuit is connected between the grid and plate, while the tap is connected to the cathode of the tube.

MULTIVIBRATOR — A two stage, resistance coupled amplifier with the output voltage fed back to the input to produce oscillations.

OSCILLATOR — An electron device for converting dc energy into ac energy.

PARASITIC OSCILLATIONS — Very high frequencies generated by inter-electrode and stray capacitances in resonance with lead inductance of the circuit.

PHASE SHIFT OSCILLATOR — An oscillator circuit which uses an RC phase shift circuit to provide the proper phase of feedback signal. This circuit is widely used as a low-frequency (audio) oscillator.

PIEZOELECTRIC EFFECT — The property of some crystals to generate a potential when a mechanical force is applied or to produce a mechanical force when a potential is applied.

SERIES FEED — In an oscillator circuit, an arrangement in which the tube plate current passes through part of the plate load.

IMPORTANT DEFINITIONS (Continued)

SHUNT FEED — In an oscillator circuit, an arrangement in which the tube plate current is carried by a circuit in parallel to the plate load.

TEMPERATURE COEFFICIENT — The fractional change of crystal frequency per unit change of temperature.

TICKLER COIL — Usually a small coil connected in series with the plate circuit of an electron tube and inductively coupled to the grid coil for the purpose of providing feedback to maintain oscillation.

TUNED-GRID OSCILLATOR — See Armstrong oscillator.

TUNED-PLATE OSCILLATOR — A form of feedback oscillator in which the tuned tank is in the plate circuit, and the feedback coupling coil is in the grid circuit.

TUNED-PLATE TUNED-GRID OSCILLATOR — (TPTG) — A form of feedback oscillator in which a tuned tank is employed in both the grid and plate circuits, and feedback is obtained through the internal grid-plate capacitance of the tube.

TRI-TET OSCILLATOR — An electron coupled crystal oscillator in which a tetrode (or pentode) is used as a triode in the oscillator circuit.

ULTRAUDION OSCILLATOR — A form of Colpitts oscillator in which the tube interelectrode capacitance forms the tapped tank circuit capacitance.

WEIN BRIDGE OSCILLATOR — An RC oscillator using a special RC circuit called the Wein bridge to obtain the proper phase feedback signal. Like the phase shift oscillator, this circuit is widely used in low frequency applications.

PRACTICE EXERCISE SOLUTIONS

1. An oscillator is used to convert direct current energy into alternating current energy.
2. No, an oscillator circuit generates an AC output signal without the aid of a signal source.
3. When energy is supplied to a tank circuit momentarily, the capacitor charges. After the source is removed, the capacitor discharges through the inductor. The inductor stores the energy and then returns it to the capacitor after it discharges. The process continues with a continuous interchange of energy between the capacitor and inductor until circuit losses absorb all the energy.
4. The frequency is inversely proportional to the value of capacitance and inductance. Increasing the value of L or C decreases frequency and vice versa.
5. False — The inductor serves as the source during these intervals. The capacitor serves as the source between intervals A-B and C-D.
6. 795 kHz — The frequency of oscillation is the same as the resonant frequency.

$$\begin{aligned}
 f &= \frac{.159}{\sqrt{LC}} = \frac{.159}{\sqrt{1 \times 10^{-10} \times 4 \times 10^{-4}}} = \frac{.159}{\sqrt{4 \times 10^{-14}}} = \frac{.159}{2 \times 10^{-7}} \\
 &= .0795 \times 10^7 = 795 \text{ kHz.}
 \end{aligned}$$

7. The waveform shown in Figure 5 is referred to as a damped wave.
8. The resistance in an LC circuit absorbs energy making it necessary to continuously supply energy to the circuit to develop a continuous ac output.
9. positive — L_1 and L_2 are phased so that the grid voltage aids the change that is occurring in plate current.
10. As C_2 discharges through L_2 , a voltage is developed across L_1 which drives the grid of V_1 negative.
11. No, with V_1 biased at cutoff, no initial pulse of plate current would occur when power is applied and oscillation would not occur.
12. When power is first applied, the bias on V_1 is zero.

PRACTICE EXERCISE SOLUTIONS (Continued)

13. Each time the grid of V_1 is driven positive, grid current flows charging C_1 to the indicated polarity. C_1 then discharges through R_1 developing the bias.
14. False — For proper operation the R_1C_1 time constant should be long so C_1 does not discharge to any extent between the pulses of grid current.
15. As the oscillation amplitude decreases, the bias decreases. As a result, the plate current pulses increase in amplitude restoring the oscillation amplitude to its former value.
16. An oscillator is a self excited amplifier in which a portion of the output energy is fed back to the input circuit.
17. An oscillator circuit employs positive feedback.
18. The difference between the circuits of Figure 8 and 10 is the location of the tank circuit. In the circuit of Figure 8 the tank circuit is in the plate circuit while in Figure 10 the tank circuit is located in the grid circuit.
19. L_1 is referred to as the tickler coil.
20. Feedback in the tuned-grid and tuned-plate oscillator is obtained inductively through a transformer arrangement.
21. A tapped tank coil arrangement is used to obtain feedback in a Hartley oscillator.
22. C_2 and R_1 in the circuit of Figure 11 develop grid leak bias.
23. No, reversing the top and the bottom ends of L_B in Figure 11 would reverse the phase of the feedback and prevent oscillation.
24. With a series-fed arrangement the tank circuit carries the dc plate current. In a shunt-fed arrangement, the tank circuit does not carry any of the dc plate current.
25. A series-fed oscillator has the disadvantage of the tank circuit carrying the dc plate current. As a result, the tank circuit is at a high dc potential with respect to ground.
26. The radio frequency choke (RFC) provides a high impedance connection between the power supply and the plate of V_1 at the operating frequency while providing a low dc resistance.

PRACTICE EXERCISE SOLUTIONS (Continued)

27. The Colpitts oscillator uses a tapped tank capacitor for feedback and the Hartley oscillator uses a tapped tank coil for feedback. Otherwise the circuits are similar.
28. Increasing the value of C_B causes the feedback voltage amplitude to decrease since X_C is inversely proportional to the value of capacitance.
29. Feedback is obtained in the TPTG oscillator through the grid-to-plate capacity of the tube.
30. False — In the TPTG circuit the plate circuit is tuned to a slightly higher frequency than the grid circuit.
31. The stability of an LC oscillator is affected by changes in electrode voltages, interelectrode capacity, tank circuit L to C ratio, the amount of energy taken from the circuit and the operating temperature.
32. In an ECO the load is coupled to the oscillator circuit through the tube's electron stream.
33. Yes, by supplying the screen (oscillator anode) through a divider connected to the plate supply, the changes in frequency caused by changes in plate voltage are cancelled by the change in screen voltage. Usually the plate to screen voltage ratio is set at 3 to 1.
34. The Clapp oscillator includes a series tuning capacitor which makes it a series tuned Colpitts.
35. L_1 and C_1 determine the frequency of oscillation.
36. The term "temperature coefficient" describes the change in frequency that occurs for a given change in temperature.
37. With a negative temperature coefficient, as temperature increases, frequency decreases and vice versa. On the other hand, with a positive temperature coefficient, a temperature increase results in a frequency increase and vice versa.
38. The TPTG circuit is similar to that of Figure 18.
39. The plate tank circuit, C_3L_2 , is tuned to a frequency slightly higher than the crystal frequency.
40. No, the tri-tet oscillator circuit of Figure 19 is similar to the Hartley oscillator.

PRACTICE EXERCISE SOLUTIONS (Continued)

41. Operating the oscillator from a separate regulated supply assures that variations in other circuits will not affect the oscillator frequency.
42. Parasitic oscillations are oscillations at a frequency other than the desired frequency.
43. The low value resistances lower the Q of the parasitic tuned circuits, thus preventing oscillations.
44. No, a multivibrator develops a nonsinusoidal output signal which is rich in harmonics.
45. The discharge time of C_2 through R_1 determines how long V_1 remains cut off.
46. Decreasing the value of the coupling capacitors increases the frequency since each tube remains cut off for a shorter period of time.
47. The output of a symmetrical oscillator consists of positive and negative pulses of equal duration. The asymmetrical circuit produces unequal positive and negative pulses.
48. The common cathode resistor R_1 provides feedback from V_2 to V_1 in Figure 22.
49. No, the grid of V_1 in Figure 22 is grounded.
50. When V_1 is conducting, its plate voltage is low and C_1 is discharging. C_1 charges when V_1 is cut off and its plate voltage is high.
51. Yes, the phase shift oscillator circuit develops a sinusoidal output signal.
52. The phase shift network consisting of C_1R_1 , C_2R_2 and C_3R_3 shifts the phase of the plate voltage 180° so the feedback is positive.
53. No, the phase shift oscillator circuit is generally used to generate low frequencies.
54. For proper operation in the Wein Bridge oscillator $R_1C_1 = R_2C_2$.
55. The Wein Bridge oscillator is used in the low or audio frequency range.
56. False — The blocking oscillator is a pulse generator and develops a nonsinusoidal output.
57. C_1 and R_1 control the frequency of oscillation in the circuit of Figure 25.
58. A pulse type signal resembling a "cash register" is available at the grid of the blocking oscillator.



ONE OF THE

BELL & HOWELL SCHOOLS

1094A

EXAMINATION
CHECK SHEET

1094A

1. D - SUPPOSE IN AN LC OSCILLATOR CIRCUIT THE VALUES OF L AND C ARE DOUBLED. WHAT HAPPENS TO THE OPERATING FREQUENCY? -- IT DECREASES BY A FACTOR OF 2.
Doubling the values of L and C decreases the frequency by a factor of $\sqrt{2 \times 2}$ or 2.

2. C - GRID LEAK BIAS IS USED IN AN LC OSCILLATOR CIRCUIT TO ACHIEVE -- SELF STARTING.
Grid leak bias permits self starting since the bias is zero at the beginning. As the oscillations build up, the bias increases until Class C operation is obtained.

3. D - THE DISTINGUISHING FEATURE OF A HARTLEY OSCILLATOR IS -- A TAPPED TANK INDUCTOR.
In the Hartley oscillator a tapped tank inductor provides the feedback.

4. D - IN A SERIES FED OSCILLATOR -- THE TANK COIL CARRIES DC PLATE CURRENT.
In a shunt fed oscillator, no dc plate current flows through the tank circuit.

5. B - THE DISTINGUISHING FEATURE OF A COLPITTS OSCILLATOR IS -- A TAPPED TANK CAPACITOR.
The Colpitts oscillator employs a tapped capacitor arrangement to obtain feedback.

6. D - IN THE TPTG OSCILLATOR -- THE PLATE TANK IS TUNED TO A HIGHER FREQUENCY THAN THE GRID TANK.
To obtain proper feedback phase, the plate tank is tuned to a frequency higher than the grid tank. In addition, the circuit oscillates at the frequency of the tank circuit with the highest Q.

7. D - WITH RESPECT TO THE CIRCUIT OF FIGURE 15 -- THE LOAD IS COUPLED TO THE OSCILLATOR THROUGH THE ELECTRON STREAM IN V_1 .
In the electron coupled oscillator, the tube's electron stream couples the oscillator to the load.

8. C - IN THE CIRCUIT OF FIGURE 20, WHEN V_1 CONDUCTS -- C_2 CHARGES.
When V_1 conducts C_1 discharges and C_2 charges.

9. A - IF THE VALUE OF R_3 IN FIGURE 22 IS INCREASED, -- V_2 REMAINS CUT OFF FOR A LONGER PERIOD OF TIME.
Increasing the value of R_3 increases the discharge time of C_1 thus holding V_2 cut off for a longer period of time.

10. B - SUPPOSE IN THE CIRCUIT OF FIGURE 25, THE PERIOD IS EQUAL TO 2 TIME CONSTANTS. IF THE VALUE OF R_1 IS 100 k Ω , WHAT VALUE OF C_1 IS REQUIRED FOR A FREQUENCY OF 10 kHz? -- 500 pF.
The period is equal to two RC time constants. One time constant is equal to $1 \times 10^{-4}/2$ or $.5 \times 10^{-4}$ sec. With R_1 equal to 1×10^5 ohms, the value of C_1 is

$$C_1 = \frac{T}{R_1} = \frac{.5 \times 10^{-4}}{1 \times 10^5} = .5 \times 10^{-9} = .0005 \mu\text{F} \text{ or } 500 \text{ pF}.$$

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1094A

Example: A railroad car is designed to travel

A	
B	
C	
D	

A. on rails. B. on a highway. C. in the air. D. under water.

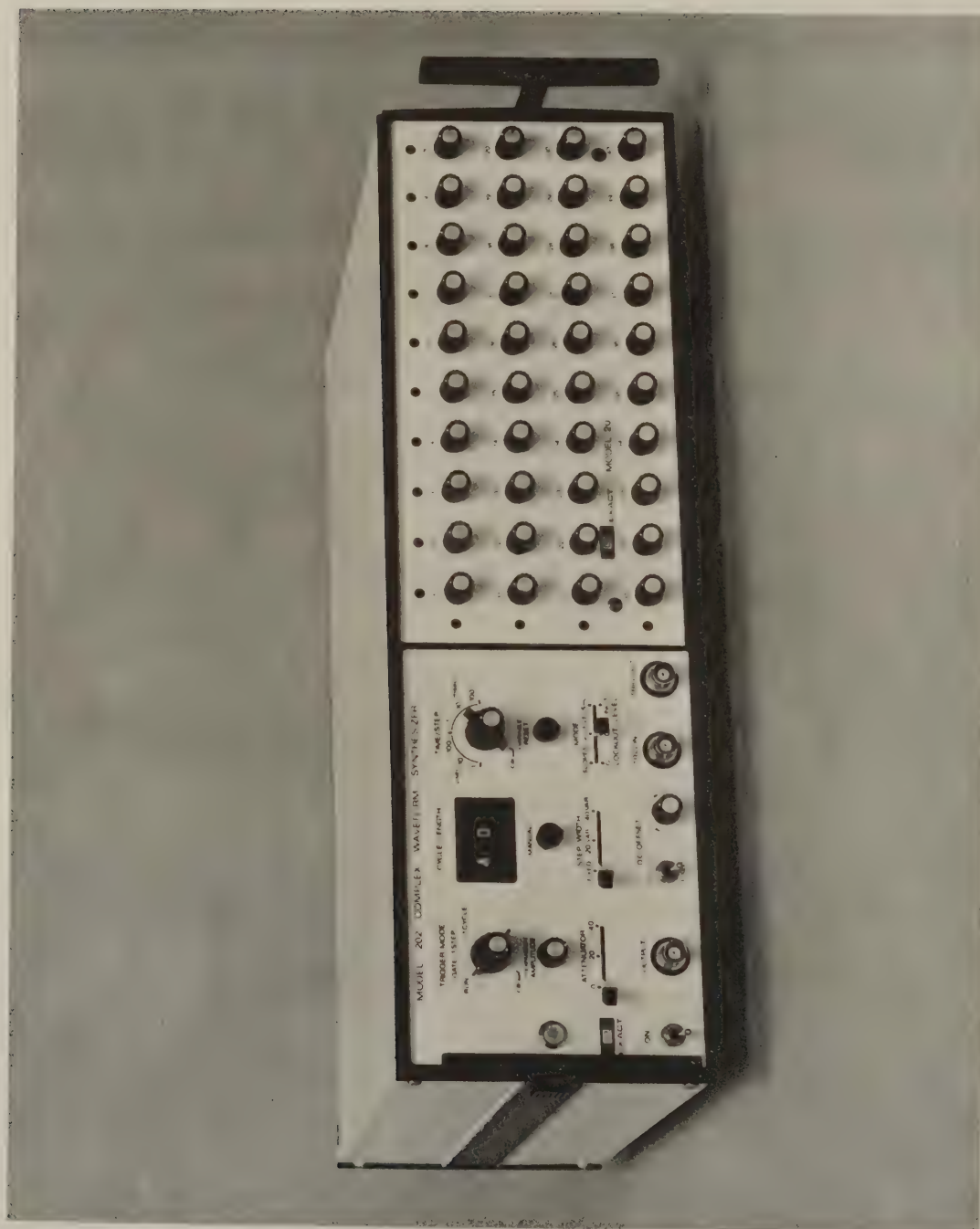
1. A ☐ Suppose in an LC oscillator circuit the values of L and C are doubled. What happens to the operating frequency?
 B ☐
 C ☒ (A) It remains the same. (B) It increases by a factor of 4. (C) It decreases by a factor of 4. (D) It decreases by a factor of 2.
 D ☐
2. A ☐ Grid leak bias is used in an LC oscillator circuit to achieve
 B ☐
 C ☒ (A) low distortion. (B) low efficiency. (C) self starting. (D) Class B operation.
 D ☐
3. A ☐ The distinguishing feature of a Hartley oscillator is
 B ☐
 C ☐ (A) a tapped tank capacitor. (B) series grid leak bias. (C) shunt feed. (D) a tapped tank inductor.
 D ☒
4. A ☐ In a series fed oscillator
 B ☐
 C ☐ (A) the plate current does not flow through any part of the tank circuit. (B) series grid leak bias must be employed. (C) a radio frequency choke is used as part of the plate load. (D) the tank coil carries dc plate current.
 D ☒
5. A ☐ The distinguishing feature of a Colpitts oscillator is
 B ☒ (A) inductive feedback. (B) a tapped tank capacitor. (C) a tapped tank inductance. (D) interelectrode capacity feedback.
 C ☐
 D ☐
6. A ☐ In the TPTG oscillator
 B ☐
 C ☐ (A) both tank circuits are tuned to the same frequency. (B) the circuit oscillates at the frequency of the tank circuit with the lowest Q. (C) the grid tank is tuned to a higher frequency than the plate tank. (D) the plate tank is tuned to a higher frequency than the grid tank.
 D ☒
7. A ☐ With respect to the circuit of Figure 15,
 B ☐
 C ☐ (A) the load is inductively coupled to the oscillator tank circuit. (B) feedback is obtained through C_{gp} . (C) C_4 and L_2 control the frequency of oscillation. (D) the load is coupled to the oscillator through the electron stream in V_1 .
 D ☒
8. A ☐ In the circuit of Figure 20, when V_1 conducts
 B ☐
 C ☒ (A) C_1 charges. (B) C_2 discharges. (C) C_2 charges. (D) the grid of V_2 is driven positive.
 D ☐
9. A ☒ If the value of R_3 in Figure 22 is increased,
 B ☐ (A) V_2 remains cut off for a longer period of time. (B) V_1 grid bias increases. (C) V_2 remains cut off for a shorter period of time. (D) V_1 conducts for a shorter period of time.
 C ☐
 D ☐
10. A ☐ Suppose in the circuit of Figure 25, the period is equal to 2 time constants. If the value of R_1 is 100 k Ω , what value of C_1 is required for a frequency of 10 kHz?
 B ☒ (A) 75 pF. (B) 500 pF. (C) 3.5 μ F. (D) 650 μ F.
 C ☐
 D ☐

TRANSISTOR OSCILLATORS

1097

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





Although sinewave oscillators are very widely used, many applications require complex wave-shape. The unit shown above, a "waveform synthesizer" can be used to develop just about any type of waveshape.

Courtesy Exact Electronics Inc.

CONTENTS

Oscillator Principles	Page 4
Oscillator Types	Page 5
LC Oscillators	Page 5
Hartley Oscillator	Page 14
Parasitic Oscillation	Page 15
Colpitts Oscillator	Page 15
Clapp Oscillator	Page 17
Frequency Stability	Page 18
Crystal Oscillators	Page 19
Relaxation Oscillators	Page 21
Blocking Oscillator	Page 22
Multivibrator	Page 24
Oscillator Output	Page 26

*The greater a man's knowledge of what has been done,
the greater is his power of knowing what to do.*

—Selected

TRANSISTOR OSCILLATORS

The basic function of an **OSCILLATOR** is to generate alternating currents or voltages. Computers use oscillators to provide synchronization and to time control operations. Oscillators are also used to generate signals in radio, television, communications, and radar transmitters and receivers. Because of this widespread use, many different types of oscillators exist. But no matter how many types of oscillators there are, they all work on the same basic principle.

OSCILLATOR PRINCIPLES

Transistor oscillators are widely used in the field of electronics. They are extremely rugged, small in size, light in weight, and require only a small amount of power for operation. This makes the transistor oscillator especially useful wherever size, weight and power requirements are a problem.

A transistor oscillator is simply an amplifier with part of its output fed back to its input. If certain conditions are met, this produces oscillating or alternating voltages or currents. This basic idea is illustrated in Figures 1 through 3.



Figure 1

Suppose a signal such as voltage V_{IN} of Figure 1 is applied to the input of an amplifier. The output of the amplifier is voltage V_O . V_O is larger in amplitude than V_{IN} and its polarity is reversed. This is normal common-emitter amplification. Let us then apply output voltage V_O to a feedback circuit as in Figure 2. The components in the feedback circuit may change the amplitude of the signal. Let's assume that these components are arranged so the signal of the feedback circuit, V_F , is equal in amplitude and phase to input voltage V_{IN} .

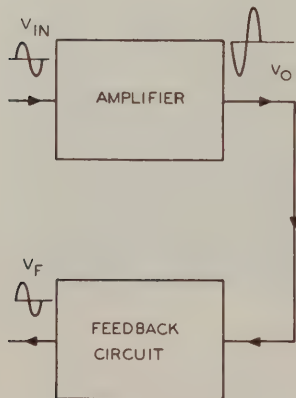


Figure 2

Notice that the polarity of the feedback voltage V_F is opposite to that of the output voltage V_O . In fact it is now the same as input voltage V_{IN} . Let us summarize this very important circuit action shown in Figure 2. V_{IN} is amplified and reversed in polarity by the amplifier, producing V_O . This is normal common-emitter amplifier action. V_O is reversed in polarity and changed in amplitude by the feedback circuit. As a result, the output of the feedback circuit (V_F) and the input signal (V_{IN}) are identical signals because they have the same amplitude and the same phase or polarity.

Now suppose we feed V_F to the input of the amplifier in place of the external input signal V_{IN} as shown in Figure 3. The feedback voltage V_F is then

amplified, producing the output voltage V_o . V_o is applied to the feedback circuit. This circuit again feeds the feedback voltage V_F back to the input circuit, and the whole operation repeats itself. The circuit uses a part of its own output for an input, which in turn is amplified to maintain the output. The circuit is now producing an alternating or oscillating output without the help of any external signal. Therefore, a circuit which operates as shown by the block diagram of Figure 3 is called an oscillator.

For the circuit to maintain the oscillations, two main requirements must be met.

1. The feedback circuit must provide a signal of proper phase to the input.
2. The feedback circuit must provide a signal of proper amplitude to the input.

Notice that the waveforms in Figures 1 through 3 are sinewaves. Sinewaves are used here to show circuit action. The waveform generated by an oscillator is determined by the components in the circuit, and may be sinusoidal or irregular in shape. The important point is this: **THE OSCILLATOR PRODUCES AN ALTERNATING OR OSCILLATING OUTPUT WITHOUT ANY EXTERNAL INPUT SIGNAL.**

We almost always want oscillations to occur at a special frequency. The job of frequency control is usually performed by the feedback circuit. The components in the feedback circuit provide a feedback voltage large enough in amplitude and of the proper phase to maintain oscillation only at the desired frequency. At all other frequencies these conditions are not satisfied.

OSCILLATOR TYPES

Many different types of **FEEDBACK OSCILLATORS** are in use today. The main difference between them is the method by which feedback is obtained. It is then logical to identify oscillators by the components used in obtaining feedback. The three most common feedback methods use: (1) an LC circuit, (2) an RC circuit, or (3) a crystal.

LC OSCILLATORS

LC oscillators use inductors and capacitors in the feedback circuit to provide feedback voltage of the correct amplitude and phase necessary for oscillation.

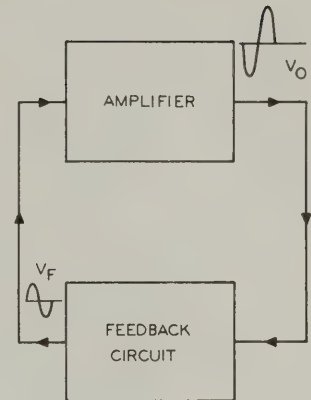


Figure 3

Figure 4 shows an LC oscillator using a PNP transistor to provide amplification. The L_1C_1 **TANK CIRCUIT** acts as the collector load for transistor Q_1 . Output voltage V_O is developed across L_1 and C_1 . Transformer coupling between coils L_1 and L_2 induces the feedback voltage into L_2 . This feedback voltage is applied to the input, the base of transistor Q_1 , to be amplified further. Thus we have an amplifier and a feedback circuit similar to Figure 3.

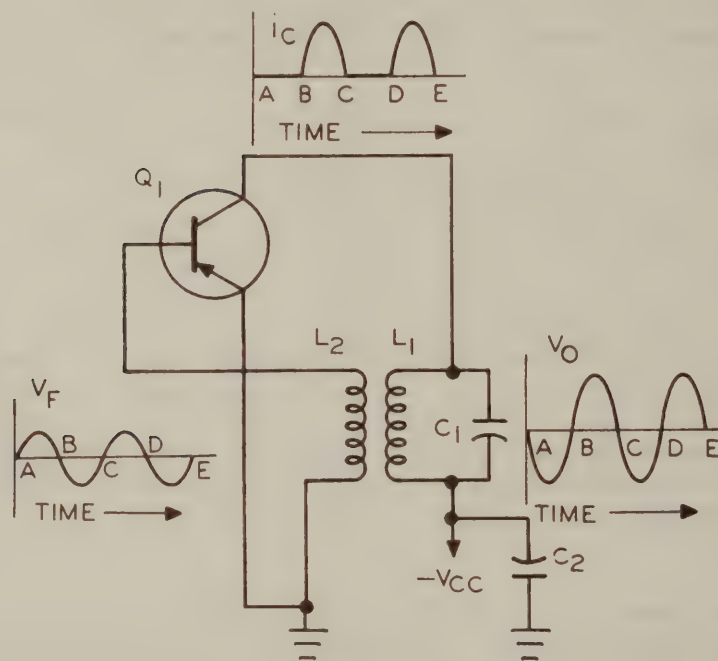


Figure 4

Also shown in Figure 4 are voltage and current waveforms of the circuit after it reaches steady operation. We can see how the circuit works by studying the relationships between these waveforms. Notice that these waveforms have a common time scale indicated by times A, B, C, D and E. That is, time A in all three waveforms, i_C , V_F , and V_O , represents exactly the same instant in time. The same is true for the corresponding points B, points C, and so on.

Assume that the circuit is already operating. The feedback voltage V_F applied to the input of PNP transistor Q_1 controls the collector current i_C of the transistor. During the time intervals from A to B and C to D, the polarity of feedback voltage V_F makes the base positive with respect to the emitter, reverse biasing the transistor. Thus, during the intervals from A to B and C to D of the i_C waveform, collector current is zero.

When voltage V_F causes the base to become negative with respect to the emitter, during intervals B to C and D to E, collector current i_c flows in the circuit. This current consists of electron flow from the negative terminal of $-V_{CC}$, through the tank circuit L_1C_1 to the collector of transistor Q_1 , and from the Q_1 emitter to ground. The current flowing through the tank circuit develops a voltage across L_1 and C_1 .

The current i_c flows only part of each cycle as shown from time B to time C and also D to E. The flywheel or ringing action of the tank circuit develops a circulating current through L_1 and C_1 on alternate half cycles as from time C to time D. This circulating tank current would gradually die out, but collector current i_c supplies energy to the tank during part of each cycle and the tank current continues to circulate. The circulating current also sets up an alternating magnetic field around L_1 . Transformer coupling between coils L_1 and L_2 develops an alternating voltage across L_2 . This is the feedback voltage V_F which is applied to the input of transistor Q_1 .

The windings of coils L_1 and L_2 in Figure 4 are connected to provide a reversal of polarity between output voltage V_o and feedback voltage V_F . Transistor Q_1 is connected as a common-emitter amplifier so the input V_F and output V_o are of opposite polarity. These two reversals of polarity produce output and feedback voltages of the proper polarity to produce oscillation. Notice that V_o and V_F of Figure 4 are the same as those of Figure 3.

The number of times each second that the tank circuit current reverses itself depends mostly on the resonant or ringing frequency of the L_1C_1 combination. This, in turn, determines the number of full cycles of voltages V_F and V_o that occur each second. This is the frequency of oscillation. However, when a transistor and other circuit components are connected to an LC tank, the various circuit values also affect the oscillation frequency. For example, the resistances of the transistor junctions as well as their capacitances have some effect. As a result, practical transistor oscillators always operate at a frequency that is somewhat lower than the resonant frequency of their tank circuit. The difference in frequency is usually small enough that for practical purposes we consider the oscillator frequency to be equal to the resonant frequency of the tank. Thus, you can find the approximate oscillation frequency of a given LC oscillator by using the resonant frequency equation:

$$f = \frac{1}{2\pi\sqrt{LC}} \text{ or } \frac{.159}{\sqrt{LC}}$$

where:

f is the frequency in hertz at which the entire circuit oscillates,

L is the inductance of the tank circuit coil in henries,

C is the capacitance of the tank circuit capacitor in farads,

π is a constant (number) approximately 3.1416.

In transistor oscillator circuits, the oscillations start at a low amplitude but build very quickly. The small initial feedback voltage variations applied to the input are amplified by the transistor. Hence, an output voltage is produced at the collector that is equal to the feedback voltage times the amplification provided by the transistor. In each of the first few cycles the pulse collector current is a little larger than in the previous cycle. A little more energy is thus given to the tank, and the oscillations increase in amplitude. The increase in oscillation amplitude continues for several cycles until the amplitude of the output reaches a level which is the same on each succeeding cycle.



A frequency counter provides a convenient means of checking the frequency of an oscillator. The time base used by the counter is generally provided by an internal oscillator. This unit employs a 1 MHz oscillator as the time base generator.

Courtesy Heath Co.

In many transistor oscillators the normal level of oscillation causes collector current to vary all the way from cutoff to saturation. Usually, zero base current is considered the cutoff point, because the collector current is nearly zero at this point. The saturation condition is that in which the base-collector junction also becomes forward biased. During a part of each cycle this junction then has very low resistance and the collector current is very high.

In some cases, the generated alternating voltage must be free of distortion. Examples are high quality oscillators used in laboratories. In Figure 4, the current i_c flows only on alternate half cycles. The flywheel or ringing action of L_1C_1 tank must provide the remaining half of voltage V_o and smooth V_o

into a sinusoidal shape. If a perfect sinewave output is required, a smaller amount of collector current flowing over the complete cycle may be an advantage. In such cases, the transistor is operated in Class A. This means that forward base bias must be provided for transistor Q_1 and a minimum amount of feedback is used. Where some distortion is permitted, oscillators are operated in Class B.

If the waveform of the output is not too important, more distortion can be tolerated to gain the advantage of Class C operation, where current i_c flows for less than half of each cycle. Shorter pulses keep the tank circuit action going and result in better efficiency. To obtain these shorter pulses of i_c , the base-emitter junction is reverse biased and operated Class C. Current i_c flows only when the base is forward biased by the peaks of base voltage.

If the transistor is operated Class A, the forward bias provides the initial pulse. However, if Class B or C operation is employed, the transistor is normally cut off. Thus, the oscillator will not start. To overcome this problem, some initial forward bias is employed. Reverse bias is then developed after oscillation starts.

Figure 5 shows an oscillator circuit which develops reverse bias after oscillation begins. The required change of bias is provided by the resistor and capacitor arrangement of R_3 and C_3 . The rest of the circuit is similar to that of Figure 4 except the feedback voltage is coupled to the emitter circuit rather than the base circuit.

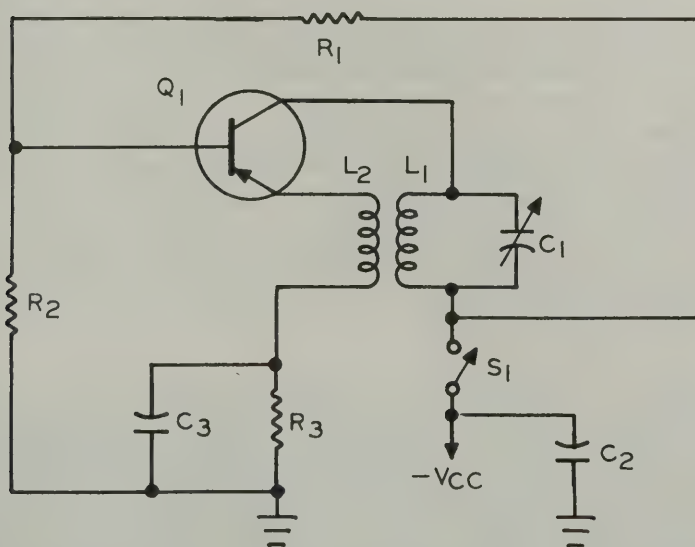


Figure 5

With switch S_1 open, collector current is zero and the circuit is at rest. Closing switch S_1 initiates circuit action. With forward bias applied to transistor Q_1 by R_1 and R_2 , collector current begins to flow. As collector current begins to increase, the magnetic field around L_1 expands. The expanding L_1 field cuts L_2 and induces a voltage into it. This voltage is positive at the top end of L_2 and negative at the lower end and thus introduces additional forward bias into the base-emitter circuit of Q_1 . This increase in forward bias causes an increase in collector current which results in further induced forward bias.

Collector current continues to increase until it reaches the saturation point. When this point is reached, the field around L_1 begins to collapse and the voltage induced in L_2 is reversed. This decreases the base bias and begins to reduce the collector current.

While collector current is flowing through R_3 , capacitor C_3 is being charged so that its upper plate is negative with respect to its lower plate. Since the base of transistor Q_1 is connected to the positive side of C_3 through R_2 , the voltage across C_3 applies reverse bias to the base-emitter circuit of Q_1 . The induced voltage across L_2 is in series with the voltage across C_3 . When the field around L_1 collapses due to a decrease in Q_1 collector current, the induced voltage reduces forward bias, and further reduces the collector current. The voltage across C_3 also acts to reduce the collector current. The induced voltage and C_3 voltage become great enough to reverse bias the base-emitter junction, thereby cutting off collector current. Even though Q_1 collector current is zero, current continues to circulate in the L_1C_1 tank circuit due to its flywheel action. This provides the opposite half of the sinusoidal voltage across L_1 and C_1 , as was shown for Figure 4.

On the peak of the next positive alternation of the flywheel current in the L_1C_1 circuit, the voltage induced in L_2 overcomes the reverse bias set up by the charge on C_3 . Thus develops a short pulse of collector current which helps induce a forward bias for the base-emitter junction, and the cycle begins again. The pulse of collector current flows through R_3 and it also replaces any of the charge that might have leaked off C_3 . Thus, Q_1 conducts for only a short period of time and Class C operation is established. Actually, a number of alternations are necessary to establish the desired voltage across C_3 . Because collector current flows for only a small part of each cycle, this class of operation is the most efficient.

In Figures 4 and 5, the tank, or tuned circuit, is in the collector circuit. The circuits of Figures 4 and 5 are then called **TUNED-COLLECTOR OSCILLATORS**. In the circuit of Figure 6, which we shall study next, the L_1C_1 tank is in the base circuit. This arrangement is called a **TUNED-BASE OSCILLATOR**. Don't forget Figures 5 and 6 are both LC oscillators. The names tuned-collector and tuned-base merely indicate the location of the tuned LC circuit.

The circuit of Figure 6 operates much like that of Figure 5. Capacitor C_3 and the various resistors perform the same functions as in Figure 5. Capacitor C_4 is a dc blocking capacitor. It prevents the low dc resistance of L_1 from shunting R_2 . This would change the voltage division between R_1 and R_2 and provide incorrect base-emitter bias.

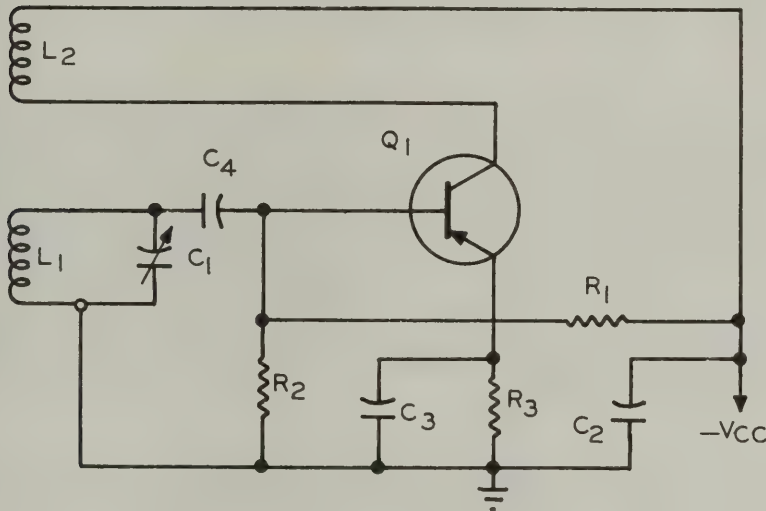


Figure 6

An additional bias voltage is provided by C_4 and R_2 . When the circulating tank current causes the upper end of the tank to become negative, base current flows from the tank through C_4 to the base of transistor Q_1 . This current charges capacitor C_4 so it is negative at the plate connected to the tank. As the flywheel tank current causes the upper end of the tank to become less negative, C_4 begins to discharge through L_1 and R_2 . The discharge current through R_2 causes a voltage across R_2 which is positive at the end connected to the base. This voltage is small at first. As oscillations increase in amplitude, the charge on capacitor C_4 increases. The voltage across resistor R_2 due to C_4 discharge also increases.

Let's pause and summarize this action. The base-emitter junction is initially forward biased by the R_1 and R_2 voltage divider. As the first several cycles of oscillation occur, an increasing voltage is developed across R_2 by the discharge of C_4 . At the same time a voltage is developed across R_3 by emitter current. All three of these voltages control the collector current of transistor Q_1 . These last two voltages are of a polarity which opposes the forward base bias applied to transistor Q_1 by the voltage divider made up of R_1 and R_2 .

In this type of circuit, components are usually selected so the reverse bias voltage developed by C_4 and R_2 is larger than the forward bias provided by the R_1R_2 divider. R_3 is often made small so very little voltage is developed across it by collector current pulses. Q_1 is then normally cut off and conducts only on the most negative peaks of the tank voltage alternation. R_3 and C_3 remain in the circuit to provide temperature stabilization. For example, if the transistor temperature increases, the collector current begins to increase. This current increase causes an increase in the reverse bias developed across R_3 , which opposes the current increase.

As explained previously, oscillators are named according to the components in the feedback circuit—LC oscillators, RC oscillators, and crystal oscillators. LC oscillators are further named according to the relative position of these feedback components. Thus, an LC oscillator can be either a tuned-collector or a tuned-base oscillator. Special features of oscillators are also used to further identify them. In the LC oscillators of Figures 4, 5, and 6, coil L_2 couples the feedback signal either to or from the L_1C_1 tank circuit. Coil L_2 is called a tickler coil. Thus, these oscillators are also called **TICKLER COIL OSCILLATORS.**

In any oscillator, sufficient ac voltage or current must be continuously fed back to the input, or the oscillations will die out. In tickler-coil oscillators, the amount of feedback depends on the amount of flux from the output circuit coil that cuts the input circuit coil. Hence, you can vary the feedback by changing the amount of coupling between the coils. One way of doing this is by moving the tickler coil with respect to the tank coil. Another way is to change the number of turns on the coils.

You may have noticed that tank capacitor C_1 is variable in the oscillator circuits of Figures 5 and 6. The oscillation frequency can be varied by adjusting C_1 . In the tuned-collector circuit of Figure 5, the dc supply voltage is applied to the tank circuit components. Thus, the tuning capacitor cannot be grounded for shielding purposes if one side of the supply is grounded, since this would short out the supply.

An advantage of the tuned-base oscillator of Figure 6 is that there is no dc voltage applied to the coil L_1 or the capacitor C_1 . This means that one side of the variable capacitor C_1 can be connected to ground. Thus, the problem of mounting the variable tuning capacitor C_1 is simplified. Also, with the frame of C_1 grounded, it helps shield the circuit from external effects which might cause the oscillator to be unstable.

Since C_1 is the means of setting the oscillator frequency, this frequency should depend on the L_1C_1 tank only. However, there is capacitance across

the tickler coil also. This is made up of internal capacitance of transistor Q_1 , capacitance between the windings of L_2 , and stray circuit capacitances. The total of all these capacitances, C_s , forms a tank circuit with coil L_2 , and this combination is resonant at some frequency. If the L_1C_1 tank losses are high at this frequency, the L_2C_s tank may take control of the oscillations. When this happens, a tuned-collector oscillator changes to a tuned-base oscillator or a tuned-base oscillator is converted to a tuned-collector oscillator. The L_1C_1 tank loses control, and it is then impossible to adjust the oscillation frequency with C_1 .

This undesired action is avoided by choosing the proper values for L_1 and L_2 to make the L_2C_s tank resonant far above the desired operating frequency of the oscillator. However, the resonant frequency of L_2C_s still limits the highest frequency at which the oscillator can operate before it becomes unstable. Usually, the upper frequency limit of the tickler-coil oscillator is about 50 megahertz.

Hartley Oscillator

If you make the tickler coil part of the tank circuit, you can avoid the unstable effects due to the tickler coil's resonant frequency. This is done in the Hartley oscillator as shown in Figure 7. Notice the tap on coil L_1 . This is the identifying feature of a **HARTLEY OSCILLATOR**. It enables coil L_1 to perform the functions of both coil L_1 and coil L_2 of the tickler coil oscillator.

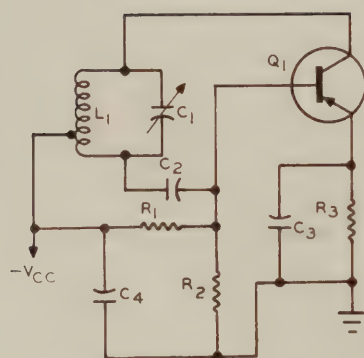


Figure 7

In Figure 7, resistors R_1 and R_2 supply forward bias to the base-emitter junction of Q_1 to start the oscillator. When dc power is applied to the circuit, the initial surge of collector current through the upper part of L_1 causes an oscillating current to be set up within the L_1C_1 tank circuit. Further collector current pulses maintain this condition so an alternating voltage is established across coil L_1 and capacitor C_1 .

The voltage developed across the part of L_1 below the tap is applied to the base of transistor Q_1 . This alternating voltage adds to and subtracts from the bias on the Q_1 base and causes a continuously pulsating collector current. These pulses of collector current supply energy to the tank and the circuit continues to oscillate. Capacitor C_2 isolates the tank from the base of Q_1 , as far as dc is concerned. Thus, the tank and the Q_1 base can be at different dc potentials, even though they are at the same ac potential.

This circuit may be biased in several ways depending on circuit values. A common practice is to depend on R_1 and R_2 to establish forward base bias to start the circuit action. R_3 is made small so very little bias is developed from collector current pulses through it. To obtain class C operation, the base current which flows on negative peaks is used to charge C_2 as explained for Figure 6. Thus, the circuit can become self biasing due to the dc voltage appearing across C_2 .

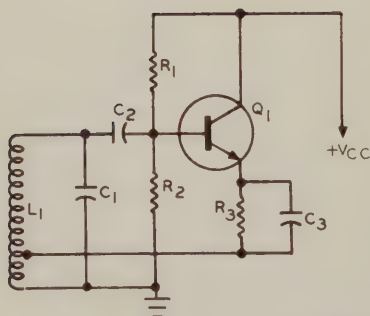


Figure 8

The Hartley oscillator of Figure 8 is very similar to that of Figure 7 except that the tap of coil L_1 is connected to the emitter circuit and Q_1 is an NPN transistor. R_1 , R_2 , R_3 , C_2 , and C_3 perform the same functions as in Figure 7. When dc power is supplied to this circuit, an initial pulse of collector current flows through the lower part of L_1 and sets up oscillating currents in the tank. In this circuit, voltage developed across the upper part of L_1 is applied to the input of Q_1 . As explained for Figure 7, this circuit may be biased in several ways depending on circuit values.

In the circuits of Figures 7 and 8, the tank inductor is a tapped coil. THIS TAPPED COIL ARRANGEMENT DISTINGUISHES THE HARTLEY OSCILLATOR FROM ALL OTHERS. The Hartley can be used at higher

The following Practice Exercise questions cover the subjects which you have just studied. They are:

OSCILLATOR PRINCIPLES

OSCILLATOR TYPES

LC OSCILLATORS

1. How is feedback employed to make an amplifier function as an oscillator?
2. To support oscillation, the feedback signal must be of the correct _____ and _____.
3. An oscillator always develops a sinusoidal output signal. True or False?
4. Does an oscillator require an input signal to develop an output?
5. Explain how a sinusoidal output voltage occurs in Figure 4 when collector current is in the form of pulses.
6. If L_1 has a value of .1 mH and C_1 has a value of 400 pF in the circuit of Figure 4, what is the approximate operating frequency?
7. What would happen in the circuit of Figure 4 if the connections to L_2 are reversed?
8. Why is some initial forward bias required in a transistor oscillator operating Class B or Class C?
9. What is an advantage of Class C operation? A disadvantage?
10. What is the function of R_1 in the circuit of Figure 5?
11. What is the function of C_3 and R_3 in the circuit of Figure 5?
12. The circuit of Figure 5 is a tuned-base amplifier. True or False?
13. What components in the circuit of Figure 6 determine the frequency of oscillation?
14. How is a reverse bias developed across R_2 in the circuit of Figure 6?
15. How can the amount of feedback be controlled in the circuit of Figure 6?
16. What factor limits the upper frequency of a tickler coil oscillator?

(

at (

ter)

frequencies than the tickler-coil types because it is not limited by the resonant frequency of the feedback section of its tank coil.

The circuits of Figures 7 and 8 are series-fed Hartley oscillators. A transistor oscillator is series fed when the collector current passes through any part of the L_1C_1 tank.

Parasitic Oscillation

Any oscillation in a circuit at an undesired frequency is called **PARASITIC OSCILLATION**. Oscillator circuits may have more than one frequency of oscillation due to the inductances of the leads and stray capacitances. These can set up small tank circuits not indicated on a schematic. Parasitic oscillations may occur at some high frequency when such undesired tank circuits exist. These oscillations absorb power from the oscillator circuit, reducing its power output at the desired frequency. They also provide unwanted signals. Careful layout of the circuit helps to avoid this.

Colpitts Oscillator

The Hartley oscillator uses a tapped coil to obtain feedback voltage. Another way to provide feedback is to divide the capacitive branch of the tank. This is done in Figure 9 by the two separate capacitors C_{1A} and C_{1B} connected in series across the coil. **THIS IS THE IDENTIFYING FEATURE OF A COLPITTS OSCILLATOR.**

In the oscillators previously described, the transistors were operated as common-emitter amplifiers. The reason we have stressed the oscillator using a common-emitter amplifier is that the majority of oscillators in use today are of this type. You will, however, find oscillators using common-base amplifiers. By rearranging the feedback components, either type of amplifier can be used in these oscillator circuits.

In the Colpitts oscillator of Figure 9, the transistor is operated as a common-base amplifier. Capacitor C_2 places the base at ac ground potential. Because the output signal is in phase with the input in a grounded base amplifier, the feedback components do not have to provide a phase reversal. The output is already of the proper phase to sustain oscillation when fed back.

In Figure 9, the voltage across C_{1B} is the feedback voltage. Although its amplitude is less than that of the tank voltage, the voltage across C_{1B} has the same polarity as the total tank voltage at any instant. The feedback is

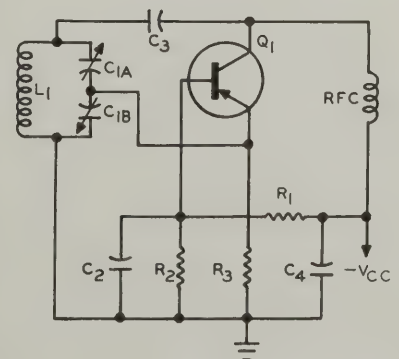


Figure 9

applied to the emitter of transistor Q_1 which is the input in this grounded-base circuit. The lower end of the tank is grounded. Hence, the feedback voltage is applied across emitter resistor R_3 . This resistor is not bypassed. A bypass capacitor across R_3 would act as a short circuit across C_{1B} at the desired oscillation frequency. C_3 blocks dc from the tank circuit. Without C_3 , the dc voltage applied to the collector would be determined by voltage division across the series dc resistance of coil L_1 and choke RFC. This could result in low collector voltage and a heavy drain on the battery.



The variable frequency oscillator (VFO) used to tune this receiver is a modification of the Colpitts oscillator circuit. Special attention is given to mechanical and electrical stability to eliminate drift.

Courtesy Heath Co.

When the oscillations cause the ungrounded end of the tank to swing positive, the feedback voltage across capacitor C_{1B} swings positive also. Since it is applied to the Q_1 emitter, this voltage change is the same as swinging the base negative. Collector current increases, the voltage across the radio-frequency choke RFC increases, and the collector voltage swings less negative. This positive swing of collector voltage is coupled by C_3 to the upper end of the tank. Notice that this positive voltage pulse is coupled to the tank at the same time that the alternating tank voltage is swinging positive. Hence, the applied pulse aids the tank voltage swing. This is another way of saying that the feedback must be of the proper phase to maintain oscillation.

The Colpitts oscillator of Figure 9 is a shunt fed oscillator. An oscillator is shunt fed when dc collector current does not flow through any part of the tank circuit. In Figure 9, C_3 prevents dc collector current from flowing through the tank. The r-f choke, RFC, offers very little opposition to dc. Thus, the dc collector current flows through the choke. The opposition of

an inductor increases as the frequency of the signal increases. At the oscillation frequency, the choke has a very large opposition to current and allows very little alternating signal to pass through it. As explained, this alternating signal is coupled by C_3 to the tank circuit to maintain oscillation.

A Colpitts oscillator is useful at the higher frequencies as it does not develop parasitic oscillations as easily as the Hartley. A disadvantage is that a circuit such as that of Figure 9 requires two variable capacitor sections, C_{1A} and C_{1B} . This increases the cost of the oscillator. However, L_1 can be made variable to provide inductance tuning. Fixed tank capacitors can then be used in place of C_{1A} and C_{1B} .

Clapp Oscillator

Figure 10 shows a modification of the Colpitts circuit called the Clapp oscillator. Except for the added capacitor C_2 in series with tank coil L_1 , this circuit is like that of the Colpitts oscillator. In this circuit, transistor Q_1 operates as a grounded-emitter amplifier. As in the Colpitts, C_{1A} and C_{1B} provide the feedback, which is applied directly to the base. A blocking capacitor is not needed at this point because there is no dc path from the base through the tank. The Q_1 collector is shunt fed through the inductor, or choke, labeled RFC. The three tank capacitors, C_{1A} , C_{1B} and C_2 , prevent direct current flow between the collector, base, or ground and the tank.

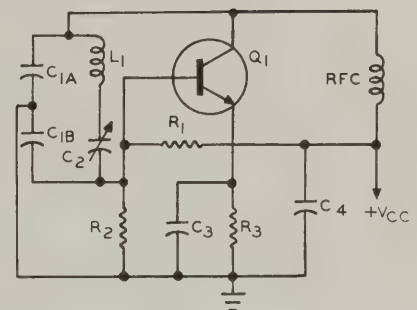


Figure 10

Basically, the oscillating action is the same as in the Colpitts oscillator. However, the operating frequency is determined by series capacitor C_2 . The capacitance of C_2 is smaller than the equivalent series capacitance of C_{1A} and C_{1B} . As a result, the oscillation frequency depends mainly on C_2 and L_1 . Circuit capacitances are shunted by the large capacitances, C_{1A} and C_{1B} , and have only a minor effect on the frequency of operation. Adjustment of C_2 selects the desired frequency.

FREQUENCY STABILITY

In many applications, frequencies must be kept very constant. This ability to oscillate at one exact frequency is known as frequency stability. No oscillator is perfectly stable. Changes of oscillator frequency are mostly due to:

(Changes in the mechanical arrangement of the parts.

(Changes in the amplification, junction resistances and capacitances in transistors.

(Changes in circuit values.

Changes in the mechanical arrangement of the circuit are caused by vibration, temperature changes, or electrostatic and electromagnetic forces. In most cases these are kept small by careful design and by temperature control.

Variations in the current gain, base and collector resistances, and interelectrode capacitances of a transistor are caused by changes in the voltages applied to the electrodes, and by changes in temperature. To prevent these changes, the operating voltages may be stabilized by voltage-regulating devices. The effect of changing interelectrode capacitance on frequency may be minimized by employing a large capacitance and small inductance in the tank and by using circuits in which the tank capacitance shunts the transistor capacitance. With this arrangement, the transistor's interelectrode capacitance is but a small part of the total capacitance. Then, any variation in the interelectrode capacitance causes only a slight frequency change.



In an FM receiver, the local oscillator may operate in the 98 to 118 MHz range. Any drift in this oscillator will result in loss of the signal. To prevent this, a special automatic frequency control (AFC) circuit is employed to keep the oscillator on frequency.

Courtesy Eico Electronic Instrument Co. Inc.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

HARTLEY OSCILLATOR

PARASITIC OSCILLATION

COLPITTS OSCILLATOR

CLAPP OSCILLATOR

17. What is the distinguishing feature of a Hartley oscillator?
18. Across what portion of L_1 is the feedback voltage developed in the circuit of Figure 7?
19. What components in the circuit of Figure 7 would develop the reverse bias voltage for Class C operation?
20. The circuit of Figure 8 is series fed. True or False?
21. What is meant by the term "parasitic oscillation"?
22. What is the identifying feature of a Colpitts oscillator?
23. In the circuit of Figure 9, Q_1 operates as a common _____ amplifier.
24. How can the amount of feedback be adjusted in the circuit of Figure 9.
25. In the circuit of Figure 9 the oscillator is (series fed) (shunt fed)?
26. What is the advantage of a Colpitts oscillator over that of a Hartley oscillator?
27. How does the Clapp oscillator differ from the Colpitts oscillator?
28. In the circuit of Figure 10, C_{1A} or C_{1B} controls the frequency of oscillation. True or False?

Changes in the circuit elements are mostly the changes in inductance, capacitance and resistance that occur as a result of temperature variations. Changes of inductance and capacitance may be kept small by careful location of component parts, by temperature control, and by the use of temperature-compensating capacitors. Such capacitors change their capacitance with temperature to offset other undesired changes in component values.

An oscillator's frequency stability depends a great deal upon the loaded Q of the tank. This loaded Q is written Q_L and is the figure of merit of the tank circuit when connected in the working circuit. The higher the Q_L , the greater the stability. Q_L is found by dividing any impedance Z_L across the tank by the reactance X of either the coil or the capacitor in the tank. This can be written $Q_L = Z_L/X$. Q_L is large, and thus the stability of the oscillator is good if X is made small.

Since X_L equals $2\pi fL$, the inductive reactance will be small if a coil with few turns or small inductance is used. Also, as X_C equals $1/(2\pi fC)$, the capacitive reactance will be small if the tank capacitor has a large capacitance. Thus, a tank that has a relatively small L and large C has a high Q_L because the ratio Z_L/X is large. Hence, the **FREQUENCY STABILITY IS GOOD WHEN AN OSCILLATOR'S TANK HAS A LOW L/C RATIO.**

CRYSTAL OSCILLATORS

One way to get good oscillator frequency stability is to use a crystal, such as quartz, to control the oscillator frequency. The crystal is simply a thin slice of quartz which has been carefully ground to a special size. When you bend or distort a properly cut crystal, electric charges appear on its opposite faces. This is known as the piezoelectric effect. If you apply a voltage to opposite faces of the crystal plate, it bends or distorts.

When you apply an alternating voltage across the crystal, it vibrates. If the frequency of the applied voltage is near the mechanical resonant frequency of the crystal, the vibrations have large amplitude. Thus, the crystal vibrates readily at only one frequency, which depends primarily on its physical size. Such a crystal can be used to control the frequency of an oscillator. The circuit oscillates at the resonant frequency of the crystal and is affected very little by component changes.

The crystal is mounted between a pair of holding plates as shown in Figure 11. The crystal in its holder has electric properties that we can represent as in Figure 12. Its mechanical vibrating characteristics are represented by resistance R , inductance L , and capacitance C_s . At the series resonant frequency

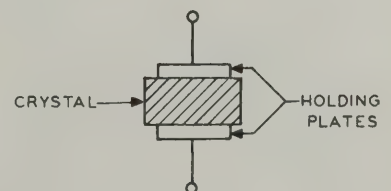


Figure 11

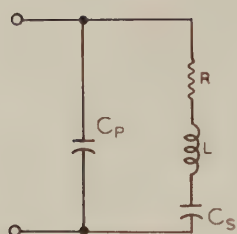


Figure 12

of L and C_s , the impedance of the crystal is at its lowest value and equal to R . This is like a series coil and capacitor arrangement except that the Q or figure of merit of the crystal is much better than that of a normal coil and capacitor. This means that at series resonance it has much lower impedance than a normal coil and capacitor.

Capacitor C_p of Figure 12 represents the capacitance between the holding plates. At frequencies higher than the series resonant frequency of L and C_s , the inductive reactance of L is greater than the capacitive reactance of C_s . Therefore, the combination of L and C_s together appear as an inductance. This inductance forms a parallel resonant tank circuit with C_p , plus any other circuit capacitance that is in parallel with C_p . At the resonant frequency of this tank, the total impedance is very high. The parallel resonant frequency depends on both the crystal and the external circuit values. By physically varying C_p and by varying external circuit values, we provide a means of making small changes in the parallel resonant frequency of the crystal.

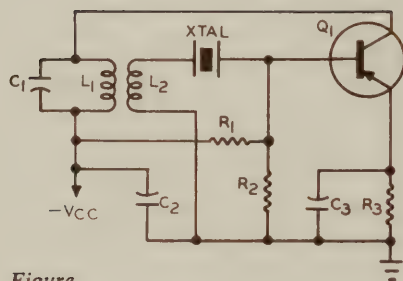


Figure 13

Figure 13 shows the circuit of a crystal oscillator. Collector tank circuit L_1C_1 is tuned to the series resonant frequency of the crystal (XTAL). Feedback to the base takes place through the inductive coupling between L_1 and tickler coil L_2 .

The crystal is in series between L_2 and the base of Q_1 . It acts as a frequency-selective filter since its impedance is minimum at its series resonant frequency. Therefore, it allows L_2 to apply maximum feedback to the input of Q_1 at this frequency. At higher and lower frequencies, the crystal has a higher impedance and reduces the amount of feedback. Therefore, the crystal prevents oscillation at frequencies other than its series resonant frequency. Other features of this circuit are the same as for the normal tuned-collector oscillator.

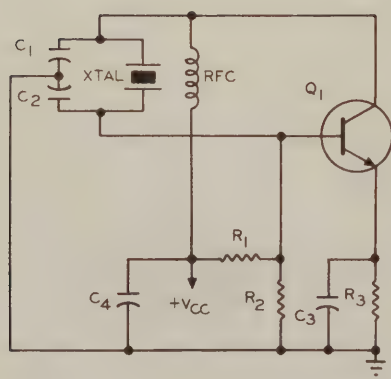


Figure 14

Figure 14 shows a Colpitts type of crystal-controlled oscillator. Here, the parallel resonant action of the crystal is used. The oscillation frequency is determined by the crystal, its holder capacitance, the equivalent series capacitance of C_1 and C_2 , and the interelectrode capacitance of transistor Q_1 . Capacitors C_1 and C_2 provide the means of tapping off the proper amount of feedback voltage from the tank for the base of Q_1 . These capacitors are small in value and are selected so that, when added to the other circuit capacitances, the total capacitance provides the desired oscillating frequency.

At parallel resonance, the impedance of the tank circuit formed by the crystal and the other circuit capacitances is maximum. Transistor Q_1 develops the greatest voltage across the tank circuit at this resonant frequency. Any variation in frequency will cause this impedance and hence tank voltage to fall rapidly.

Part of the voltage developed across the tank by circulating currents is the feedback voltage. Thus, at the resonant frequency of the tank, with the tank impedance maximum the feedback voltage will be maximum. At any other frequency the feedback voltage across C_2 is too small to maintain oscillation. Transistor Q_1 operates as a grounded-emitter amplifier in Figure 14. Thus, the junction between capacitors C_1 and C_2 is grounded so that the feedback voltage coupled from the lower end of the tank is 180° out-of-phase with the collector voltage at the upper end of the tank.

The same type of oscillator circuit is shown in Figure 15, but transistor Q_1 here operates as a grounded-base amplifier. Hence, to provide the needed in-phase feedback, the emitter is fed from the junction between C_1 and C_2 . With this arrangement, the voltage across C_2 , which is the feedback voltage, has the same phase as the collector voltage at the upper end of the tank.

Why is the crystal oscillator more stable than an LC oscillator? As you have seen they both operate in much the same manner. The difference is in the Q of the circuits. A typical value of Q for a good LC tank circuit is 200. In comparison, a typical Q for a good crystal is 50,000. Figure 16 shows the graph of impedance versus frequency for an LC tank and Figure 17 shows a graph of impedance versus frequency for a crystal. Notice that the curve representing crystal impedance is much sharper than the curve of LC tank impedance. This is the result of the much larger Q of the crystal. Thus, with equal changes in frequency, the changes in the impedance of the crystal are much larger. This is the primary reason why the crystal oscillator is more stable than the LC oscillator.

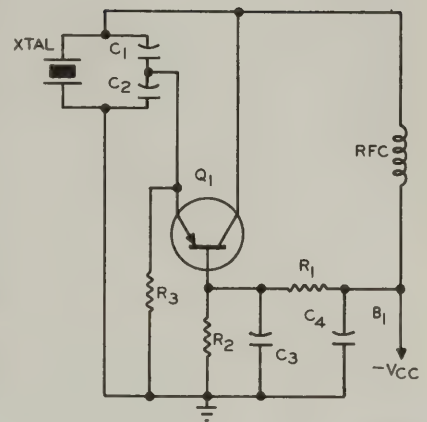


Figure 15

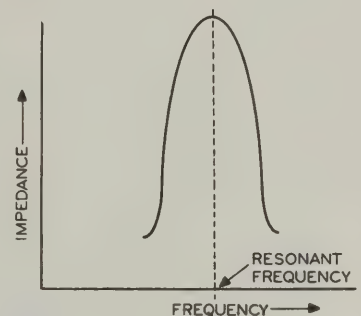


Figure 16

RELAXATION OSCILLATORS

A **RELAXATION OSCILLATOR** is generally described as an oscillator in which the frequency is controlled by the time required for a capacitor to charge or discharge through a resistor. This takes the place of the tank circuit effect of LC oscillators. In some cases the frequency may depend upon more than one RC circuit. To a lesser extent, an inductor and resistor are used to control the frequency in a relaxation oscillator. This use is so limited that relaxation oscillators are almost always described from the standpoint of RC control.

Two of the most common types of relaxation oscillators are the blocking oscillator and the multivibrator. Both types are used in many practical applications. Computers, television transmitters and receivers, and radar systems all use large numbers of relaxation oscillators.

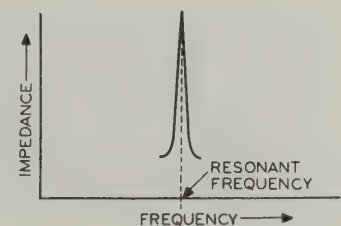


Figure 17



Blocking oscillators and multivibrators are often used in TV receivers to generate the horizontal and vertical sweep signals.

Courtesy Sony Corporation of America

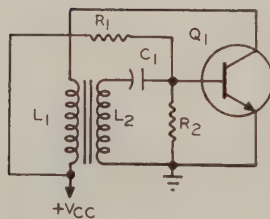


Figure 18

Blocking Oscillator

Figure 18 shows the circuit of a transistor **BLOCKING OSCILLATOR**. The collector connects through L_1 to the positive terminal of collector supply V_{CC} . The voltage divider arrangement of R_1 and R_2 provides forward bias for Q_1 .

At the instant the supply circuit is closed, there is a rise of forward current in the base of Q_1 . An amplified current is produced in the collector circuit.

This current through L_1 induces a voltage across L_2 . By properly phasing the windings, the induced voltage at the ungrounded end of coil L_2 is now positive. The voltage across inductor L_2 charges capacitor C_1 through the normally small base-emitter resistance to make the right-hand plate of C_1 negative with respect to the left-hand plate.

The charging current through the capacitor is carried by the base-emitter junction. This current adds to the original forward base current. This larger base current is further amplified in the collector circuit. This process of current buildup in the collector circuit continues rapidly until the transistor is saturated.

When collector saturation is reached, the collector current, and hence the current through L_1 , becomes constant. Since the current through L_1 is no longer changing, there is no longer inductive coupling from L_1 to L_2 and the induced voltage across L_2 drops to zero. As a result, there is no longer any charging voltage across C_1 and it starts to discharge through R_2 . The ungrounded end of R_2 is negative with respect to the grounded emitter, and a reverse bias is applied to the base-emitter circuit.

With reverse bias applied and no induced voltage present across L_2 , current in the base decreases toward zero and, in turn, the collector current decreases toward zero. This decrease of collector current in L_1 induces a voltage in L_2 which makes its ungrounded end negative. This voltage rapidly reduces the forward base current and collector current to zero.

Capacitor C_1 continues to discharge, but more slowly, depending upon the values of R_2 and C_1 . The voltage across R_2 , due to this discharge current, decreases to a value where it can no longer hold Q_1 cut off. Then forward bias current starts and the cycle repeats. Notice that **THERE IS NO RESONANT CIRCUIT ACTION HERE.** When the transistor is cut off it remains that way until C_1 discharges sufficiently through R_2 . How soon the cycle repeats depends on the speed of this discharge. Thus, the frequency of oscillation is again controlled by part of the feedback circuit, but now it is an RC combination. The blocking oscillator frequency of Figure 18 can be varied by changing R_2 , C_1 or both. Although the transistor in Figure 18 is an NPN type, a PNP type could be used. Then, the polarity of bias voltage V_{CC} must be reversed.

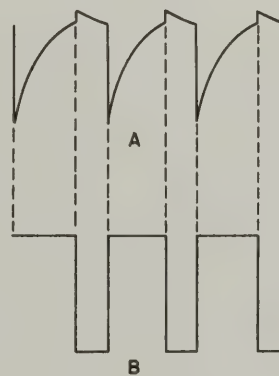


Figure
19

The blocking oscillator does not develop a sinusoidal output signal. Figure 19A shows the waveshape appearing at the base of Q_1 and Figure 19B shows the waveshape appearing at the collector. The exponential rise in the base waveshape occurs when C_1 in Figure 18 discharges. Since Q_1 goes alternately from cutoff to saturation, the collector waveshape has a rectangular shape.

Multivibrator

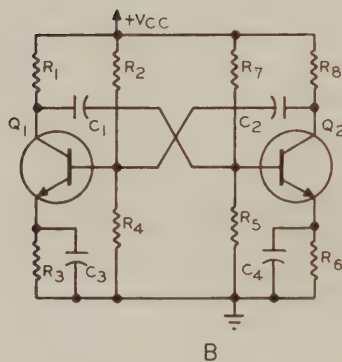
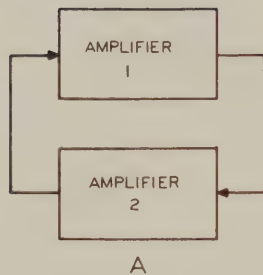


Figure 20

The basic collector-coupled transistor **MULTIVIBRATOR** of Figure 20 is a two-stage resistance-capacitance coupled common-emitter amplifier. The output of the first stage is coupled to the input of the second stage and the output of the second stage is coupled to the input of the first stage. Figure 20A shows a block diagram of this circuit. Compare it to Figure 3. Notice that the feedback circuit of Figure 3 has been replaced by amplifier 2. In this type of oscillator, the feedback loop includes an amplifying transistor rather than a tank circuit or transformer.

With this circuit arrangement, the current in one transistor must be increasing while the current in the other is decreasing. This continues until a point is reached at which the transistors reverse their conditions. Thus, there is a sort of see-saw action where the transistors alternately conduct and cut off. Since there is no steady or stable state, the output is an alternating or oscillating voltage.

To follow circuit operation, let's assume that the circuit has been in operation for some time. Assume capacitor C_1 has a positive charge at the plate connected to the collector of Q_1 . Assume also that the collector current through transistor Q_1 is increasing.

As the collector current increases and the voltage across R_1 increases, the collector voltage decreases (becomes less positive). Therefore capacitor C_1 begins to discharge. The discharge path is through R_5 , the parallel combination of R_3 and C_3 and transistor Q_1 . The resultant voltage across R_5 reduces the forward bias of the emitter-base junction of Q_2 . That is, the base of Q_2 is made less positive. Thus, the base current and collector current of Q_2 decrease. The collector voltage of Q_2 increases (becomes more positive) and C_2 begins to charge to this new value through the parallel circuit made up of R_4 and the base-emitter resistance of Q_1 . The resultant voltage across R_4 further increases the forward bias on Q_1 , and Q_1 conducts even heavier than before. This increase in conduction of Q_1 and decrease in conduction of Q_2 continues until transistor Q_1 is saturated and transistor Q_2 is cut off.

Q_2 remains cut off until the discharge of C_1 begins to slow down and the voltage across R_5 begins to decrease. As this voltage decreases, the reverse bias at the base of Q_2 also decreases. Q_2 then begins to conduct and the voltage at the collector of Q_2 begins to decrease (becomes less positive). Capacitor C_2 discharges through R_4 , the parallel combination of C_4 and R_6 , and Q_2 .

As this action continues, the reverse bias voltage applied to the emitter-base junction of Q_1 increases and Q_1 collector current decreases. Thus the col-

lector voltage of Q_1 increases (becomes more positive) and C_1 begins to charge to this value through R_5 . The resultant voltage across R_5 increases the forward bias on Q_2 , which further increases conduction. Due to the feedback arrangement of the circuit, the increased collector current of Q_2 continues to decrease the collector current of Q_1 . The decrease on Q_1 collector current in turn makes Q_2 conduct still more. This action continues until Q_2 is saturated and Q_1 is cut off.

The circuit remains in this state until the voltage developed across R_4 by the discharge of C_2 can no longer hold Q_1 in cutoff. Q_1 begins to conduct, collector current increases, and the collector voltage decreases. Notice that we have arrived at the same conditions as those which were assumed when we began the description. The complete cycle repeats itself over and over, thus providing an oscillating condition.

The current in one transistor is increasing as that in the other is decreasing. These two conditions are constantly reversing between the two transistors. THIS IS CHARACTERISTIC OF MULTIVIBRATORS. The frequency of oscillation depends primarily on the values of R_5 , C_1 , R_4 and C_2 . As in the blocking oscillator, capacitor discharge through resistance determines the cutoff period and hence the length of each cycle.

As in the case of the blocking oscillator, the signal developed by a multivibrator is not a sinusoidal signal. The signal at the bases of Q_1 and Q_2 in Figure 20 is similar to that shown in Figure 19A. The signals at the collectors are rectangular like that shown in Figure 19B.

While many types of multivibrators are in use, the basic action described for Figure 20 is common to all. Multivibrators employ feedback voltages between two amplifier stages that are usually large enough to cut off the stages alternately. Output voltage is usually taken from one of the collectors. The main difference in multivibrators is the manner in which feedback voltage is obtained and applied. The knowledge of the basic operation of the collector-coupled multivibrator will aid you in the understanding of any multivibrator you may encounter.

A characteristic of relaxation oscillators is that their output frequency is not very stable. They are affected by temperature changes, variations in the supply voltage, and conditions at the transistor junctions. Where accuracy in timing is needed, as in computer, radar and television applications, a triggering or synchronizing voltage is applied to the oscillator from an external source, usually a crystal-controlled unit. This triggering voltage can be used to bring the cutoff transistor into conduction at the desired instant. This causes a certain portion of the oscillator cycle to occur at regular time intervals and provides stable oscillation frequency.

OSCILLATOR OUTPUT

The method of obtaining oscillator output varies widely with the type of oscillator and the particular application. In the figures of this lesson, we have omitted the output coupling elements for simplicity, as such elements are not part of the oscillator circuit proper.

There are two types of coupling which you will encounter most often. In the first, the signal is capacitively coupled from either the input or output circuits. In the second, the output is inductively coupled from the oscillator tank circuit. The inductance of the tank circuit acts like the primary of a transformer when another winding is added close to it and the signal is taken from the secondary. As these oscillator circuits are found in various types of equipment, the method of output coupling is easily seen.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

FREQUENCY STABILITY

CRYSTAL OSCILLATORS

RELAXATION OSCILLATORS

BLOCKING OSCILLATOR

MULTIVIBRATOR

OSCILLATOR OUTPUT

29. Name some factors which can affect the frequency stability of an LC oscillator.
30. For good stability, an oscillator should have a high L to C ratio. True or False?
31. What factor determines the resonant frequency of a crystal?
32. Can a crystal act as only a series resonant circuit?
33. In the circuit of Figure 13 the crystal operates in its parallel mode. True or False?
34. How can the frequency of oscillation be varied in the circuit of Figure 14?
35. Q_1 in Figure 15 operates as a common-emitter amplifier. True or False?
36. Why is a crystal oscillator more stable than an LC oscillator?
37. In the circuit of Figure 18, as Q_1 collector current increases the base is driven _____.
38. Increasing the value of R_2 in Figure 18 causes the operating frequency to decrease. True or False?
39. What type of waveshape is available at the collector of Q_1 in Figure 18?
40. What is the name given to the circuit shown in Figure 20B?
41. In the circuit of Figure 20B, are Q_1 and Q_2 conducting at the same time?
42. What happens to the operating frequency in the circuit of Figure 20B if the values of C_1 and C_2 are decreased?
43. In the circuit of Figure 20B, if Q_1 is conducting C_2 is discharging. True or False?

(1)

(2)

(3)

SUMMARY

An oscillator is an electronic generator of alternating voltage or current. An LC feedback oscillator is an amplifier in which part of the output is coupled back to the input to produce oscillation. The frequency of oscillation is determined by the values of inductance and capacitance in the tank circuit. The feedback must have enough amplitude and must be of the proper phase to sustain oscillation. Those circuits that use a separate feedback coupling are called tickler-coil oscillators. Feedback may also be obtained by tapping the tank coil as in a Hartley oscillator or tapping the tank capacitance as in a Colpitts oscillator.

To obtain greater frequency stability, crystal control of the oscillator may be employed. In these circuits the frequency depends primarily on the physical size of the crystal which acts like a very high Q tank circuit.

A relaxation oscillator is an oscillator in which the feedback circuit includes resistance and capacitance in place of inductance and capacitance. The frequency of oscillation is controlled by the time required for one or more capacitors to charge or discharge through one or more resistors. Two common types of relaxation oscillators are the blocking oscillator and the multivibrator. In the blocking oscillator feedback is taken from a transformer and then applied to an RC circuit. The multivibrator uses a second transistor to obtain phase inversion and so there are then two RC combinations to influence the frequency. Relaxation oscillators are often controlled by signals from more stable sources, such as crystal oscillators.

IMPORTANT DEFINITIONS

BLOCKING OSCILLATOR — An ac generator in which the feedback is provided through induction, but the oscillation frequency is determined mainly by the RC values in the input circuit.

CLAPP OSCILLATOR — A Colpitts type of oscillator in which a small capacitor is included in the inductive branch of the tank to improve frequency stability.

COLPITTS OSCILLATOR — A feedback oscillator in which the feedback is obtained by tapping the capacitive branch of the tank circuit.

FEEDBACK OSCILLATOR — An amplifier in which energy fed back from the output to the input causes oscillations.

HARTLEY OSCILLATOR — A feedback oscillator in which the feedback is obtained by tapping the inductive branch of the tank circuit.

MULTIVIBRATOR — A two-stage resistance-capacitance coupled amplifier with the output voltage fed back to the input to produce oscillations.

OSCILLATOR — An electron device for generating ac.

PARASITIC OSCILLATION — An undesired oscillation in a circuit.

RELAXATION OSCILLATOR — Generally an oscillator in which the frequency is controlled by the time required for a capacitor to charge or discharge through a resistor. In some cases the frequency may depend upon more than one RC circuit. Also an oscillator in which the frequency is controlled by an LR circuit.

SERIES FED — In an oscillator circuit an arrangement in which the dc collector current passes through any part of the tank circuit.

SHUNT FED — In an oscillator circuit an arrangement in which the dc collector current is carried by a circuit in parallel with the tank circuit.

TANK CIRCUIT — A parallel LC circuit. In an oscillator, the oscillation frequency is determined by the resonant frequency of the tank. Also referred to as a tuned circuit.

TICKLER-COIL OSCILLATOR — Any feedback oscillator in which feedback is obtained by means of a separate tickler coil that is inductively coupled to the tank coil.

TUNED-BASE OSCILLATOR — A feedback oscillator in which the tank is in the base circuit. Feedback is provided by a tickler coil.

TUNED-COLLECTOR OSCILLATOR — A feedback oscillator in which the tank is in the collector circuit. Feedback is provided by a tickler coil.

PRACTICE EXERCISE SOLUTIONS

1. To make an amplifier function as an oscillator, the output signal is fed back to the input.
2. phase, amplitude — To sustain oscillation the feedback signal must have the correct phase and amplitude.
3. False — The output signal waveshape depends upon the components in the circuit.
4. No, an oscillator does not require an input signal. The input signal is provided by the feedback circuit.
5. The resonant circuit "flywheel effect" converts the pulses of collector current into a sinusoidal signal.
6. The frequency of oscillation is approximately the resonant frequency of the L_1C_1 tank circuit.

$$\begin{aligned}
 f &= \frac{.159}{\sqrt{LC}} = \frac{.159}{\sqrt{1 \times 10^{-4} \times 4 \times 10^{-10}}} = \frac{.159}{\sqrt{4 \times 10^{-14}}} \\
 &= \frac{.159}{2 \times 10^{-7}}
 \end{aligned}$$

$$f = .0795 \times 10^7 = .795 \times 10^6 = .795 \text{ MHz or } 795 \text{ kHz}$$

7. If the connections to L_2 are reversed, the phase of the feedback signal is reversed, and the circuit will not oscillate.
8. Without the initial forward bias, the circuit would not be self starting. The forward bias provides the initial pulse of current to start the oscillations.
9. Class C operation has the advantage of high efficiency but the disadvantage of increased distortion.
10. R_1 provides forward bias for Q_1 so the circuit will start oscillating when power is applied.
11. C_3 and R_3 develop reverse bias for Class C operation in the circuit of Figure 5.
12. False — The circuit of Figure 5 is a tuned-collector amplifier since the tuned circuit is in the collector circuit.
13. L_1 and C_1 determine the frequency of oscillation.

PRACTICE EXERCISE SOLUTIONS (Continued)

14. When the base of Q_1 is driven negative, base current charges C_4 . As the signal drives the base of Q_1 positive, C_4 discharges through R_2 developing the reverse bias.
15. The amount of feedback can be controlled by varying the coupling between L_1 and L_2 or the number of turns on L_1 and L_2 .
16. The upper frequency of a tickler coil oscillator is limited by the stray capacity which resonates with the tickler coil.
17. The distinguishing feature of a Hartley oscillator is the tapped tank coil.
18. The feedback voltage is developed across the lower portion of L_1 .
19. C_2 and R_2 would develop the primary reverse bias for Class C operation. The reverse bias provided by R_3 and C_3 would be for temperature stabilization.
20. True — collector current passes through the lower portion of L_1 .
21. Parasitic oscillations are oscillations that occur at frequencies other than the desired frequencies.
22. The identifying feature of a Colpitts oscillator is a tapped tank capacitance to provide the feedback voltage.
23. base — C_2 places the base of Q_1 at signal ground thus allowing it to operate as a common-base amplifier.
24. The amount of feedback is controlled by varying the values of C_{1A} and C_{1B} .
25. shunt fed — The circuit of Figure 9 is shunt fed since no dc collector current flows through the tank circuit.
26. The Colpitts oscillator is more stable with regard to parasitic oscillations than the Hartley.
27. The Clapp oscillator is a series tuned Colpitts oscillator.
28. False — C_2 controls the frequency of oscillation. C_{1A} and C_{1B} have only a minor effect on the operating frequency.
29. The stability of an LC oscillator is affected by changes in the mechanical arrangement of the parts, changes in component values and changes in transistor gain, resistance and capacity.

PRACTICE EXERCISE SOLUTIONS (Continued)

- 30. False — For good stability an oscillator should have a low L to C ratio so stray circuit capacity is a small part of the total capacity.
- 31. The resonant frequency of a crystal is determined primarily by its physical size and shape.
- 32. No, a crystal can act as both a series and a parallel resonant circuit.
- 33. False — The crystal is in series with the feedback circuit and operates in its series resonant mode. At resonance, the crystal impedance is minimum and feedback is maximum.
- 34. The parallel resonant frequency can be varied by changing the value of C_1 and C_2 .
- 35. False — Q_1 operates as a common-base amplifier.
- 36. The Q of a crystal is greater than that of a tank circuit and likewise the response curve for a crystal is sharper than that of an LC tank circuit.
- 37. positive — Since Q_1 is an NPN transistor, the base must be driven positive to increase collector current. L_2 is phased so that this situation occurs.
- 38. True — Increasing the value of R_2 increases the discharge time of C_1 thus decreasing the frequency.
- 39. A rectangular waveshape, like that shown in Figure 19B, is available at the collector of Q_1 .
- 40. The circuit of Figure 20B is called a multivibrator.
- 41. No, Q_1 and Q_2 conduct alternately in the multivibrator circuit.
- 42. Decreasing the values of C_1 and C_2 decreases their discharge time, increasing frequency.
- 43. False — When Q_1 is conducting, Q_2 is cut off and C_2 is charging.



1097A

1097A
EXAMINATION
CHECK SHEET

1. D - AN ELECTRONIC OSCILLATOR CIRCUIT -- EMPLOYS POSITIVE FEEDBACK.
Positive feedback is employed in an oscillator since the feedback must aid the input signal.
2. C - IF R_1 IS REMOVED FROM THE CIRCUIT OF FIGURE 5 -- THE CIRCUIT WILL NOT START.
Removing R_1 removes the forward bias on Q_1 . With no forward bias, Q_1 operates at cutoff and will not conduct when power is applied.
3. D - THE CIRCUIT OF FIGURE 6 IS A -- TUNED-BASE AMPLIFIER.
The circuit of Figure 6 is a tickler coil oscillator and since the tuned tank circuit is in the base it is called a tuned-base amplifier.
4. D - THE CIRCUIT OF FIGURE 7 IS -- SERIES FED.
Collector current flows through the upper part of the tank circuit and thus the circuit is series fed.
5. B - IN THE CIRCUIT OF FIGURE 10 -- THE VALUE OF C_2 IS LESS THAN THAT OF C_{1A} AND C_{1B} .
The value of C_2 is less than that of C_{1A} and C_{1B} . Therefore, C_2 mainly controls the frequency of oscillation.
6. D - WITH RESPECT TO THE CIRCUIT OF FIGURE 13 -- THE CRYSTAL OPERATES AS A SERIES RESONANT CIRCUIT.
The crystal operates as a series resonant circuit since it is in series with the feedback path to the base of Q_1 .
7. D - CONNECTING A SMALL VALUE CAPACITOR IN PARALLEL WITH THE CRYSTAL IN FIGURE 14 WILL -- DECREASE THE FREQUENCY.
The capacitor lowers the resonant frequency of the crystal since frequency is inversely proportional to capacity.
8. C - ALL OF THE FOLLOWING LC COMBINATIONS ARE RESONANT AT APPROXIMATELY 1 MHz.
WHICH WOULD GIVE THE GREATEST FREQUENCY STABILITY? -- 256 pF AND 100 μ H.
Of the choices given, this combination provides the lowest L/C ratio which provides the greatest stability.
9. A - TO INCREASE THE FREQUENCY OF OSCILLATION IN FIGURE 18 -- DECREASE THE VALUE OF C_1 .
Reducing the value of C_1 decreases the C_1 discharge line thus increasing frequency.
10. B - IN THE CIRCUIT OF FIGURE 20 -- C_1 DISCHARGES WHEN Q_2 IS CUT OFF.
 C_1 charges when Q_1 is cut off and Q_2 is conducting, and discharges when Q_1 is conducting and Q_2 is cut off.

1097A

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1097A

Example: Books are placed

A	☐
B	☒
C	☐
D	☐

A. in a bathtub. B. on a bookshelf. C. on a road.
D. in a thimble.

- 1. An electronic oscillator circuit**

A	☐
B	☐
C	☐
D	☒

(A) requires an AC input signal. (B) does not employ any feedback. (C) employs negative feedback. (D) employs positive feedback.
- 2. If R_1 is removed from the circuit of Figure 5**

A	☐
B	☐
C	☒
D	☐

(A) the frequency will increase. (B) the frequency will decrease. (C) the circuit will not start. (D) the output amplitude will increase.
- 3. The circuit of Figure 6 is a**

A	☐
B	☐
C	☐
D	☒

(A) tuned-collector amplifier. (B) Hartley oscillator. (C) Colpitts oscillator. (D) tuned-base oscillator.
- 4. The circuit of Figure 7 is**

A	☐
B	☐
C	☐
D	☒

(A) shunt fed. (B) a Colpitts oscillator. (C) Clapp oscillator. (D) series fed.
- 5. In the circuit of Figure 10**

A	☐
B	☒
C	☐
D	☐

(A) C_{1A} and C_{1B} mainly control the frequency of oscillation. (B) the value of C_2 is less than that of C_{1A} and C_{1B} . (C) R_3 and C_3 provide forward bias. (D) R_1 provides reverse bias.
- 6. With respect to the circuit of Figure 13**

A	☐
B	☐
C	☐
D	☒

(A) C_1 and L_1 control the frequency of oscillation. (B) the crystal operates as a parallel resonant circuit. (C) R_1 and R_2 provide reverse bias. (D) the crystal operates as a series resonant circuit.
- 7. Connecting a small value capacitor in parallel with the crystal in Figure 14 will**

A	☐
B	☐
C	☐
D	☒

(A) increase the frequency. (B) stop oscillation. (C) change the output voltage waveshape. (D) decrease the frequency.
- 8. All of the following LC combinations are resonant at approximately 1 MHz. Which would give the greatest frequency stability?**

A	☐
B	☐
C	☒
D	☐

(A) 10 pF and 2.56 mH. (B) 100 pF and 256 μ H. (C) 256 pF and 100 μ H. (D) 1 pF and 25.6 mH.
- 9. To increase the frequency of oscillation in Figure 18**

A	☒
B	☐
C	☐
D	☐

(A) decrease the value of C_1 . (B) increase the supply voltage. (C) increase the value of C_1 . (D) increase the value of R_2 .
- 10. In the circuit of Figure 20**

A	☐
B	☒
C	☐
D	☐

(A) C_2 charges when Q_2 conducts. (B) C_1 discharges when Q_2 is cut off. (C) reducing the value of C_1 and C_2 reduces the frequency. (D) C_2 discharges when Q_1 conducts.

FIELD EFFECT TRANSISTORS

1100

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





Many modern, high-quality television receivers use field effect transistors in the tuner section. The high input impedance of the FET prevents loading of the tuned circuit and increases selectivity.

Courtesy Radio Corporation of America

CONTENTS

Field Effect Transistors	Page 4
FET Characteristics	Page 8
Insulated Gate FET	Page 13
Operating Modes	Page 15
FET Amplifiers	Page 16
Common Source	Page 16
Common Drain	Page 17
Common Gate	Page 18
Self Bias	Page 20
Multi-Stage Amplifiers	Page 20
Direct-Coupled Amplifiers	Page 21
R-F Amplifiers	Page 23

*Let no feeling of discouragement prey upon you,
and in the end you are sure to succeed.*

—A. Lincoln

FIELD EFFECT TRANSISTORS

The conventional junction transistor has earned a permanent place in modern electronic technology. The transistor's small size, low cost and durability have brought about developments that were almost impossible with vacuum tube technology. Of course, the transistor has not completely replaced vacuum tubes. There are still many applications for vacuum tubes in modern electronic systems.

The conventional junction transistor is basically a low input impedance device. There are many applications where a very high impedance is needed. Remember, a voltmeter's internal impedance must be many times greater than the circuit impedance to prevent excessive loading. Analog computer operational amplifiers and many types of instrument amplifiers also require high input impedances.

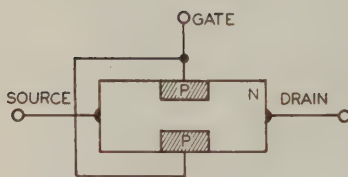


Figure
1

In some cases, high input impedances can be obtained with junction transistors by employing special circuit techniques. Vacuum tubes can also be used in applications where conventional junction transistors do not provide the desired high input impedance. In the past, this often led to the development of "hybrid" circuits employing both tubes and transistors. However, a later development in semiconductor technology, the **FIELD EFFECT TRANSISTOR (FET)**, is a semiconductor device that ~~presents a high input impedance without the need for special circuitry.~~

FIELD EFFECT TRANSISTORS

In the conventional junction transistor, the input circuit (base-emitter junction) consists of a forward biased PN junction. The impedance across a forward biased PN junction is very low (just like a conducting diode), which accounts for its low input impedance. The impedance across a reverse biased PN junction is very high (just like a cutoff diode).

The action by which current control takes place in an FET is different from that in a conventional junction transistor. Current flow in an FET is controlled by an electric field (hence the name Field Effect Transistor), which makes the control action similar to that in a vacuum tube. The FET's similarity to a vacuum tube has brought about the name "solid state vacuum tube".

The basic construction of a field effect transistor is shown in Figure 1. The device of Figure 1 starts out as a bar of silicon which behaves much like a resistor. The silicon bar can be either P-type or N-type material.

Terminals are connected to each end of the silicon bar. These are the terminals labeled **SOURCE** and **DRAIN** in Figure 1. The source and drain connections are merely ohmic contacts—no PN junctions exist at these points.

The device is completed by diffusing P-type material into both sides of the N-type bar to form PN junctions. The two diffused P-type regions are connected together and referred to as the **GATE**.

Figure 2 shows how an N-type bar material FET is connected to external power sources for proper operation. The drain is connected through R_L to the positive end of the drain supply, V_{DD} . The source is connected to the negative terminal of V_{DD} . Drain current I_D flows through the FET from source to drain and through the external circuit as shown by the arrows. Ignoring the action of the gate for the moment, drain current is limited by R_L and the resistance of the N-type material.

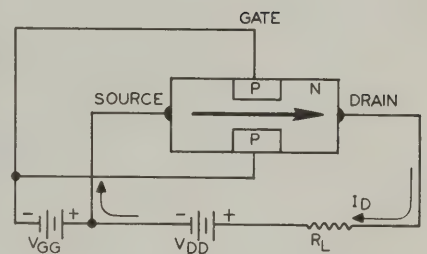
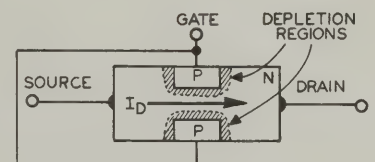


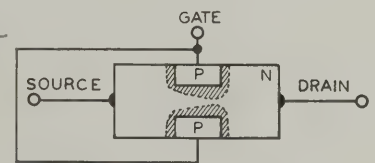
Figure
2

The gate is connected to the negative terminal of the gate supply, V_{GG} . The positive terminal of V_{GG} is returned to the source. This connection reverse biases the gate-source PN junction. With the gate reverse biased, the only gate current flowing is an extremely small reverse current.

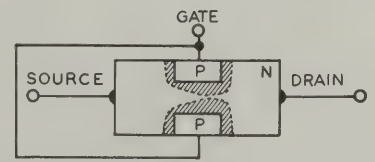
For the moment, let's assume the gate voltage is zero. As drain current I_D flows through the FET, it reverse biases the gate PN junctions. Whenever a PN junction is reverse biased, a depletion region, void of current carriers, is set up around the junction. This is shown in Figure 3A. The drain current cannot flow in the depletion regions, since no current carriers are available. Thus, drain current is confined in the space between the depletion regions. The space between the depletion regions is referred to as the CHANNEL.



A



B



C

Figure
3

As the drain-source voltage is increased, the drain current increases. However, the increase in drain current with drain-source voltage is not linear. This is because the reverse bias on the gate PN junctions increases as the drain current increases. This in turn causes an increase in the size of the depletion regions, causing them to extend further into the channel as shown in Figure 3B. Thus, as the drain current increases, it causes a decrease in channel width, which tends to oppose an increase in drain current.

Increasing the drain-source voltage further causes the depletion regions to almost come together as shown in Figure 3C. When this condition is present, further increase in drain-source voltage causes no further increase in drain current. Thus, drain current is at its saturated value.

The point where drain current reaches its saturated value is referred to as **CHANNEL PINCH-OFF**. If the drain-source voltage is increased too far, breakdown of the device occurs and it conducts heavily.

With zero gate-source voltage, pinch-off or drain current saturation is a function of drain current. Control of the point where channel pinch-off and likewise drain current saturation occurs can be obtained by making the gate negative with respect to the source. As the negative gate voltage is increased, the depletion regions are extended further into the channel. This causes channel pinch-off to occur at a lower drain-source voltage, and thus causes a lower drain current.

Therefore, drain current can be controlled by gate voltage. Making the gate more negative decreases drain current, and vice versa.

Drain current flows through N-type material in the device shown in Figure 3. This type of device is therefore referred to as an **N-CHANNEL FET**. It is possible to construct a **P-CHANNEL FET** by diffusing gates of N-type material into a block of P-type material.

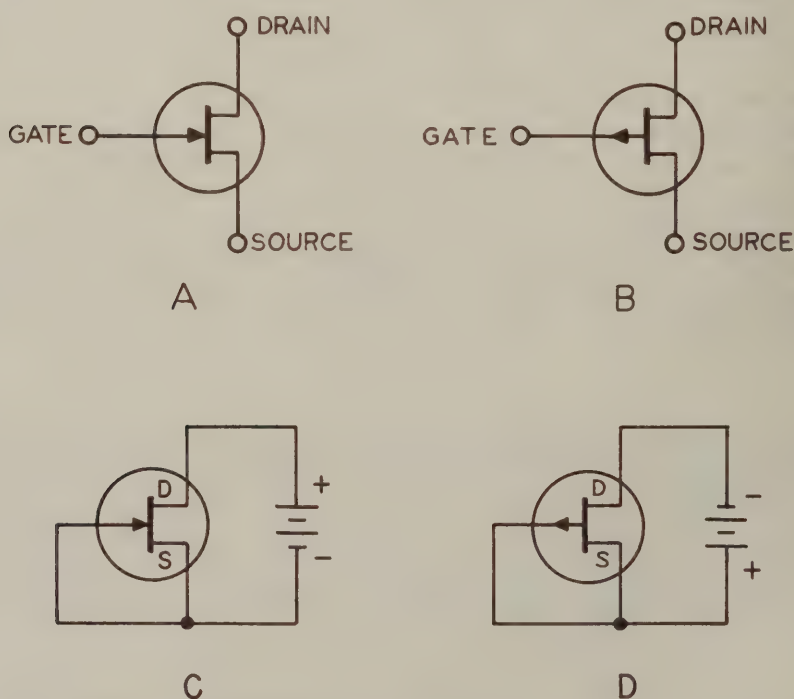
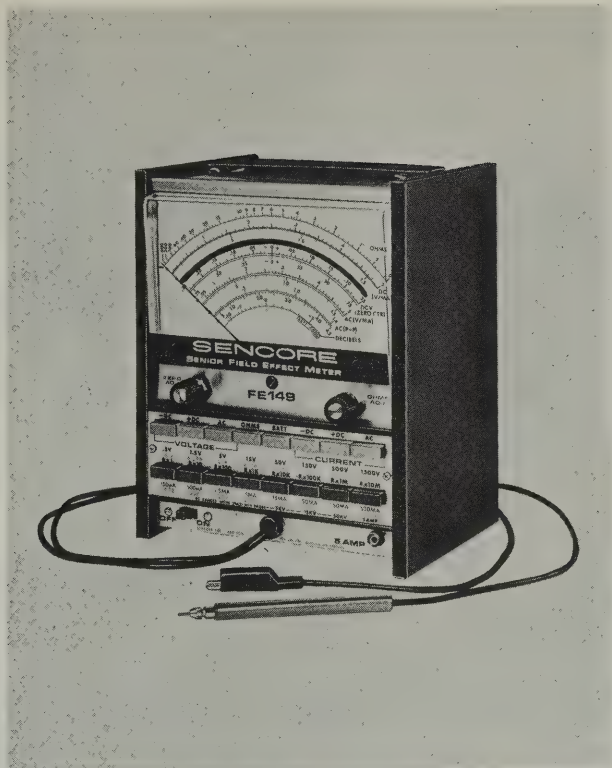


Figure 4



The VOM shown above uses field effect transistor techniques to provide a high input impedance.

Courtesy Sencore, Inc.

Figure 4 shows the symbols commonly used for field effect transistors. FET's may be either P-channel or N-channel types. Figure 4A shows the symbol for an N-channel FET and Figure 4B shows the symbol for a P-channel FET. In using the two types, it is necessary to observe polarity. An N-channel device is connected to a battery as a vacuum tube would be connected. The source is connected to the negative side, and the drain is connected to the positive side, as shown in Figure 4C. This configuration coincides with vacuum tube cathode and plate connections, and under these conditions the N-channel FET may be used as an AMPLIFIER. A P-channel device requires the reverse of these polarities, as shown in Figure 4D, with the drain connected to the negative side of the battery. Connected in this way, a P-channel FET may also be used as an amplifier. By reversing these polarities, a P- or N-channel FET may be used as a source FOLLOWER, similar to a cathode follower vacuum tube stage. A simple way to remember the drain polarity is to note the direction of the arrow on the gate. In a junction-type FET, the arrow points inward to a positive drain, and away from a negative drain. The source and drain should always be indicated on circuit diagrams by the letters S and D.

The elements of an FET can be compared to those of a conventional junction transistor and vacuum tube. The function of the gate is comparable to that of the transistor base or tube control grid. The function of the source is comparable to that of the transistor emitter or tube cathode. The function of the drain is comparable to that of the transistor collector or tube plate.

FIELD EFFECT TRANSISTORS

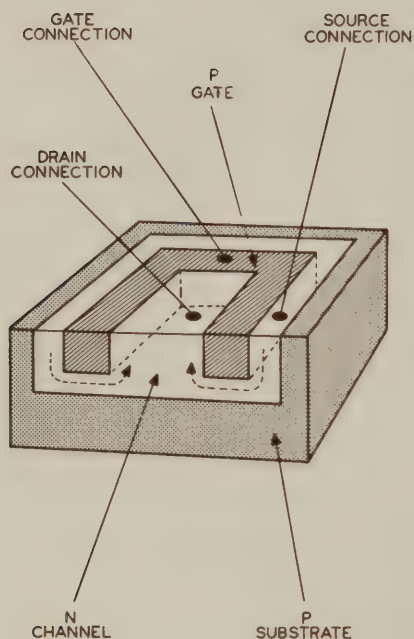


Figure
5

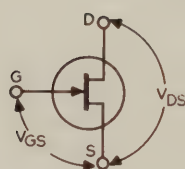


Figure
6

Diffusing the gate regions into both sides of a block of semiconductor material like that shown in Figure 3 is a difficult job. A simpler approach is to use the SINGLE-ENDED structure shown in Figure 5. To make an N-channel device, it starts out as a block of P-type material which forms the base or substrate. An N-type region is then diffused into the P-type substrate as shown to form the channel. The device is then completed by diffusing P-type gate material into the N-type channel. Connections are then made to the various materials to complete the device. Connections are made to the top surface.

FET CHARACTERISTICS

As with conventional junction transistors and vacuum tubes, special symbols are employed to designate the various electrode voltages and currents of an FET. Figure 6 shows the important electrode voltages. The drain-source voltage is labeled V_{DS} and the gate-source voltage is labeled V_{GS} .

In FET notation, special attention must be given to the subscript "S". An "S" as the first or second subscript refers to the SOURCE terminal on the device. However, when triple notation is used, an "S" as the third subscript stands for "SHORTED". This means that any terminal (or terminals) not represented by one of the first two subscripts is shorted to the common terminal, which is always the second subscript. For example, the notation I_{GSS} (defined later) represents the gate-source current with the drain tied to the source.

Figure 7 shows a set of drain characteristic curves for a typical N-channel FET. These curves show how drain current I_D varies with drain-source voltage V_{DS} for given values of gate-source voltage V_{GS} . The drain characteristic curves are similar to the collector characteristics of a conventional junction transistor and the plate characteristics of a pentode vacuum tube.

Let's follow the curve labeled $V_{GS} = 0$ in Figure 7. Many FET specifications are given by manufacturers at $V_{GS} = 0$. Reading from left to right, as V_{DS} increases, I_D increases. The curve begins to level off at a drain-source voltage of about +5 volts. Further increases in V_{DS} cause only a small increase in I_D . The value of drain current at the leveling off point is labeled I_{DSS} , drain saturation current.

Even though the gate junction is reverse biased, some small gate current flows. At the I_D saturation point, the value of gate current is referred to as the gate-source leakage current and is labeled I_{GSS} .

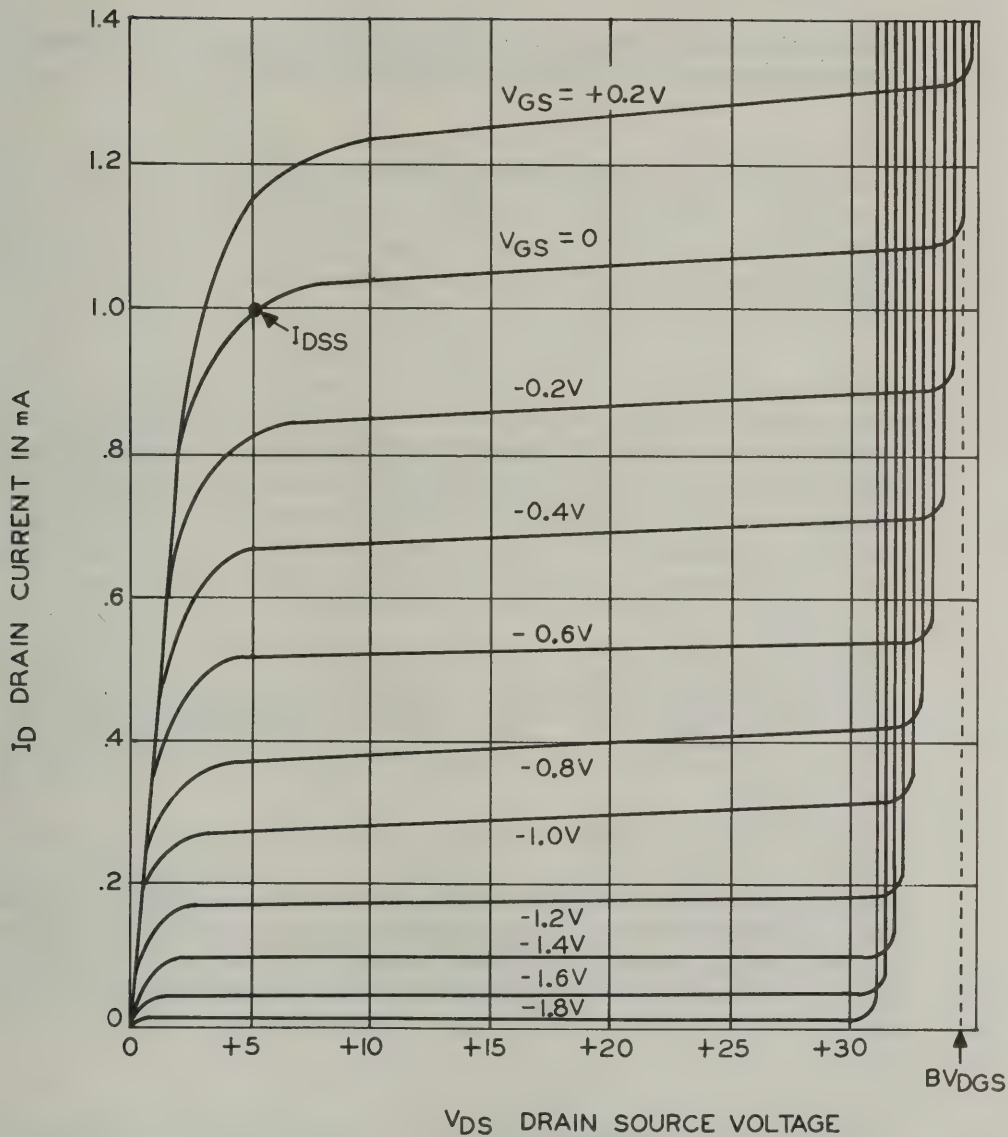


Figure 7

If the drain-source voltage is increased too far, breakdown occurs. This causes a large increase in drain current as shown by the almost vertical rise in the curve at the right. The actual device breakdown occurs between the drain and the gate. Thus, the breakdown voltage is labeled BV_{DGS} . However, the drain-gate voltage at breakdown is approximately equal to the drain-source voltage at this point. This is why the rating BV_{DGS} is shown on the drain-source voltage scale on the curves of Figure 7.

As the gate is made negative, drain current I_D decreases. This is shown by the curves in Figure 7 becoming lower as the gate-source voltage V_{GS} becomes more negative. If V_{GS} is made negative enough, I_D becomes almost zero. This

value of gate-source voltage is called the pinch-off voltage (V_p). The pinch-off voltage is comparable to cutoff voltage in a vacuum tube.

The characteristic curves shown in Figure 7 are for an N-channel FET. The curves would be similar for a P-channel FET except the polarity of the gate-source voltage V_{GS} and drain-source voltage V_{DS} would be reversed. That is, V_{GS} would be positive and V_{DS} would be negative.

The dynamic characteristics of a device are those which specify the device under actual operating conditions. For example, with vacuum tubes the dynamic characteristics are μ (amplification factor), g_m (transconductance) and r_p (internal impedance). Current gain, h_{fe} , is a typical dynamic characteristic for a junction transistor. Like these devices, the field effect transistor has its own particular set of dynamic characteristics. These characteristics enable us to determine such things as the voltage gain of an amplifier.

As mentioned earlier, the field effect transistor is a constant current device as is a pentode vacuum tube. That is, beyond a certain point (channel pinch-off), a further increase in drain voltage causes almost no further increase in drain current.

Drain current is then controlled entirely by gate voltage. Therefore, one of the most important dynamic characteristics of an FET is how much change in drain current is produced by a given change in gate voltage. This dynamic characteristic is referred to as the **FORWARD TRANSFER ADMITTANCE** (y_{fs}).

The forward transfer admittance is equal to a change in drain current divided by the change in gate-source voltage that caused the drain current change. The measurement is made while holding drain-source voltage constant.

$$\left(y_{fs} = \frac{\Delta I_D}{\Delta V_{GS}} \text{ with } V_{DS} \text{ held constant} \right) \quad (1)$$

As shown by Equation 1, y_{fs} is calculated by dividing a current by a voltage. The Δ means "a change of". This is the opposite of a resistance or impedance calculation, where a voltage is divided by a current. The opposite or reciprocal of impedance is $1/Z$ or admittance Y . Admittance, rather than conductance ($G = 1/R$), is used in the specification since it takes into account any reactive effects. The international unit of admittance is the siemens (S). However, the most widely used unit is the mho. A siemens or mho is equal to 1 ampere divided by 1 volt.

As an example, suppose a certain FET has a change in I_D of .35 mA when the change in V_{GS} is .4 volts. The forward transfer admittance is therefore

$$y_{fs} = \frac{.35 \times 10^{-3} \text{ amp}}{.4 \text{ volt}} = .875 \times 10^{-3} \text{ mho}$$

The unit mho is generally too large for most y_{fs} specifications. Generally, y_{fs} is specified in μmhos (10^{-6} mhos) or in $\mu\text{siemens}$ (10^{-6} S). Thus, $.875 \times 10^{-3}$ mhos (or siemens) becomes 875 μmhos (or $\mu\text{siemens}$).

The y_{fs} specification for an FET is similar to the transconductance (g_m) for a vacuum tube. Because of this, some FET manufacturers use the symbol g_m instead of y_{fs} to represent forward transfer admittance.

In summary, y_{fs} is a figure of merit for an FET. The higher the value of y_{fs} , the greater the amplifying ability of the FET. Generally, the y_{fs} specification is given at a signal frequency of 1 kHz. FET's intended for high frequency operation usually specify an additional value of y_{fs} at or near the desired operating frequency.

Another useful dynamic FET characteristic is the **OUTPUT ADMITTANCE (y_{os})**. The output admittance is found by dividing a change in drain current by the change in drain-source voltage that caused the drain current change. The measurement is made while holding gate-source voltage constant.

$$y_{os} = \frac{\Delta I_D}{\Delta V_{DS}} \text{ with } V_{GS} \text{ held constant} \quad (2)$$

Like forward transfer admittance y_{fs} , output admittance y_{os} is commonly expressed in micromhos.

Some FET manufacturers may specify output impedance, which is the reciprocal of y_{os} or $Z_o = 1/y_{os}$. Since admittance is the reciprocal of impedance, a low admittance means a high impedance, and vice versa.

As an example of the y_{os} calculation, suppose a particular FET has a change in I_D of .125 mA when V_{DS} is changed 5 volts. The output admittance is

$$\begin{aligned} y_{os} &= \frac{.125 \times 10^{-3} \text{ amperes}}{5 \text{ volts}} \\ &= .025 \times 10^{-3} \text{ mhos or } 25 \text{ micromhos} \end{aligned}$$

The equivalent output impedance Z_o is the reciprocal of the output admittance.

$$Z_o = \frac{1}{.025 \times 10^{-3} \text{ mhos}} = 40 \times 10^3 \text{ ohms or } 40 \text{ k}\Omega$$

The Z_o or $1/y_{os}$ specification for an FET is equivalent to the dynamic plate resistance (r_p) specification for a vacuum tube. The Z_o or y_{os} specification along with the forward transfer admittance can be used to determine the voltage gain of an FET amplifier stage.

Another useful FET characteristic is the **AMPLIFICATION FACTOR (μ)**. This is equivalent to the μ specification for a vacuum tube. It represents the theoretical maximum voltage gain that can be obtained from the FET. Amplification factor can be determined by dividing a change in drain-source voltage by the change in gate-source voltage that caused the change in drain-source voltage. The measurement is made with drain current held constant. Expressed as an equation:

$$\left(\mu = \frac{\Delta V_{DS}}{\Delta V_{GS}} \text{ with } I_D \text{ held constant} \right) \quad (3)$$

In many cases, μ is not given on an FET specification sheet. However, it can be calculated easily by using y_{fs} and y_{os} :

$$\left(\mu = y_{fs}/y_{os} \right) \quad (4)$$

As an example, let's take the values of y_{fs} and y_{os} determined earlier and calculate μ :

$$\begin{aligned} \mu &= \frac{.875 \times 10^{-3} \text{ mhos}}{.025 \times 10^{-3} \text{ mhos}} \\ \mu &= 35 \end{aligned}$$

As mentioned earlier, the dynamic characteristics of an FET are similar to those of a vacuum tube. The characteristics of both are listed below to illustrate the similarity:

FET	VACUUM TUBE
y_{fs} — forward transfer admittance	g_m — transconductance
$1/y_{os}$ — output impedance	r_p — dynamic plate resistance
μ — amplification factor	μ — amplification factor

The dynamic characteristics of an FET provide a means of evaluating circuit operation as is the case with vacuum tube dynamic characteristics.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

FIELD EFFECT TRANSISTORS

FET CHARACTERISTICS

1. The FET is a device with a (a) high input impedance, (b) low input impedance.
2. The basic operation of an FET is similar to that of a conventional junction transistor. True or False?
3. Name the three terminals of a field effect transistor.
4. PN junctions are formed at the drain and source connections. True or False?
5. The gate PN junctions of an FET are (a) forward biased, (b) reverse biased.
6. What is the channel of an FET?
7. What is channel pinch-off?
8. Increasing the negative gate-source voltage causes pinch-off to occur at a (a) lower value of drain current, (b) higher value of drain current.
9. Making the gate more negative with respect to the source causes drain current to (a) increase, (b) decrease.
10. Draw the symbols for the N-channel and P-channel FET.
11. The function of the gate of an FET is comparable to that of the _____ of a conventional junction transistor.
12. The FET structure shown in Figure 5 is referred to as (a) single-ended, (b) double-ended.

13. Define the following:

a. V_{GS}

c. I_{DSS}

e. I_D

b. V_{DS}

d. BV_{DGS}

i. I_{GSS}

14. The point at which a further increase in V_{DS} causes very little increase in I_D is called _____.

15. Define pinch-off voltage V_P .

16. V_P is similar to what vacuum tube characteristic?

17. What is y_{fs} and how is it determined?

18. The unit most widely used for y_{fs} specifications is the (a) siemens or mho, (b) millisiemens or millimho, (c) microsiemens or micromho.

19. The y_{fs} specification is similar to what vacuum tube specification?

20. What is y_{os} and how is it determined?

21. If the y_{os} of an FET is given as $.02 \times 10^{-3}$ mhos, what is the output impedance of the FET?

22. The $1/y_{os}$ specification is similar to what vacuum tube specification?

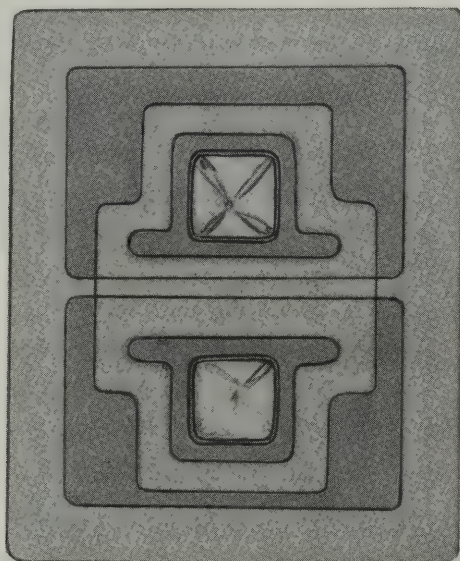
23. What is μ and how is it found?

24. How is μ determined if y_{fs} and y_{os} are known?

INSULATED GATE FET

The field effect transistors described previously employ PN junctions at the gate. This type of FET is often called a **JUNCTION FIELD EFFECT TRANSISTOR (JFET)**. The high input impedance of a JFET is due to reverse bias at the gate PN junction.

It is possible to build an FET whose input impedance is very high regardless of gate voltage polarity. This type of FET does not employ a regular PN junction at the gate. Actually, the gate is insulated from the rest of the device by an oxide or nitride layer. This type of FET is called an **INSULATED GATE FIELD EFFECT TRANSISTOR (IGFET)**.



This photomicrograph shows the physical appearance of a MFE2093 N-Channel JFET on a silicon chip.

Courtesy Motorola Semiconductor Products, Inc.

As shown in Figure 8, the insulating layer forms a small high-quality capacitor with the gate as one plate and the rest of the device as the other plate. In this way, current in the channel is controlled without the disadvantages of a gate diode. However, in order to control channel current by using an electric field, the insulating layer must be very thin. As such, the capacitor has a very low breakdown voltage — approximately 30V. The gate capacitance is so low, and the leakage resistance is so high, that the gate is easily charged to voltages above breakdown. Because of this, IGFET's must be handled with care. They should always be handled by the case, and not by the leads, since static electricity discharging through the leads during handling may be sufficient to damage the device. A grounded soldering iron is also required when installing the device in a circuit. IGFET's are generally shipped from the supplier with

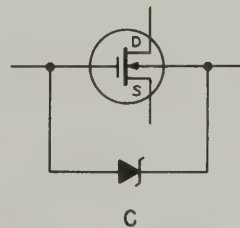


Figure 8

the leads shorted together and grounded by using a shorting ring. Also, some devices are manufactured with built-in zener diodes to protect the gate from low-energy sources. The symbol for protected IGFET's is shown in Figure 8. The zener diode is effective in protecting the device, but reasonable caution is still required in the application of insulated gate devices.

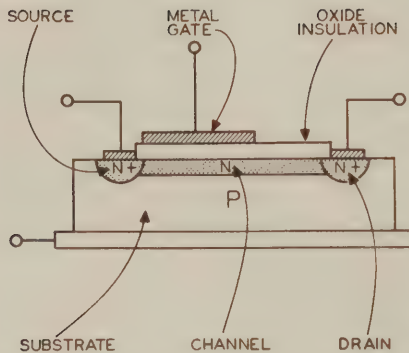


Figure 9

Figure 9 shows the structure of one type of IGFET. Heavily doped N-type regions (indicated by the N+) are diffused into a P-type substrate or base. A channel of regular N material is diffused between the heavily doped N-type regions. The source and drain connections are made to the heavily doped N-type regions.

The oxide or nitride insulating layer is then formed over the channel, and a metal gate layer is deposited over the insulating oxide layer. There is no electrical connection between the gate and the rest of the device. Note again that the gate and channel act as the plates of a capacitor, with the insulation as the dielectric.

The operation of the IGFET is basically the same as that of the JFET explained earlier. The potential applied to the gate sets up an electric field in the channel. This electric field produces a depletion region in the channel, and thus controls drain current.

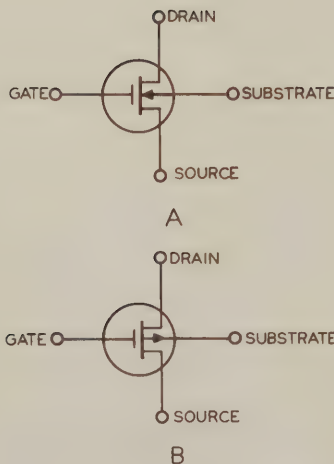
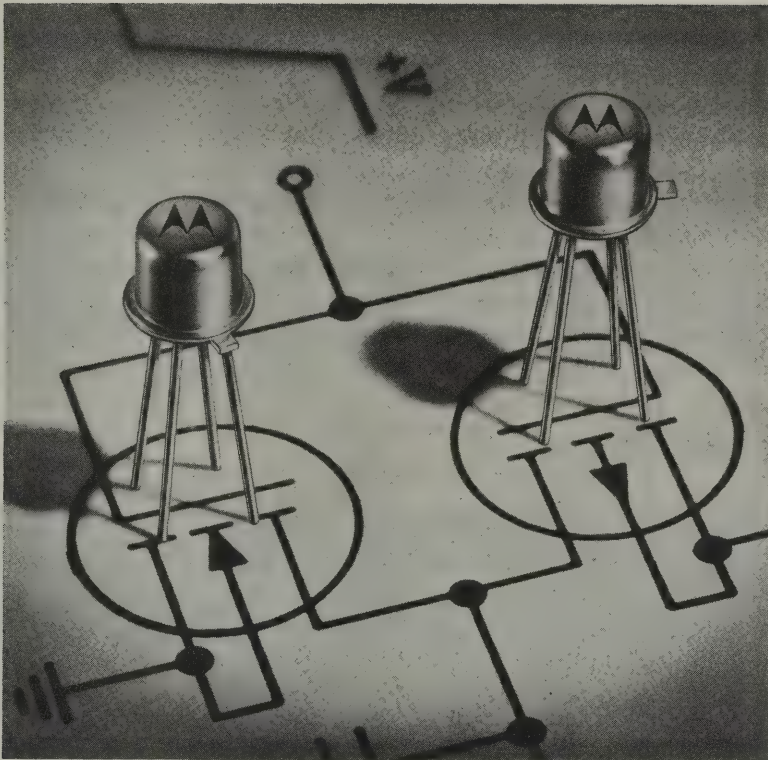


Figure 10

The oxide insulating layer formed over the channel is usually silicon oxide. Since the deposited gate material is a metal, the IGFET is often referred to as a METAL-OXIDE-SEMICONDUCTOR (MOS).

There are several symbols in use for the IGFET or MOS. Figure 10 shows one of the commonly used symbols in its N-channel and P-channel versions. Notice that the gate is shown as a capacitor to the rest of the device, since this is actually the way the device is constructed. Notice also that there is a connection to the substrate. The arrow on the substrate indicates that there is a PN junction between the channel and substrate. Generally, the substrate is connected to the source or some other point so this PN junction is zero or reverse biased.

The symbol shown in Figure 10A is for the N-channel structure shown in Figure 9. The symbol for a P-channel IGFET, as shown in Figure 10B, is the same as that for the N-channel IGFET except the arrow on the substrate is reversed.



The MOS field effect transistors shown above are designed for switching applications.

Courtesy Motorola Semiconductor Products, Inc.

OPERATING MODES

The JFET's described earlier are said to operate in the **DEPLETION MODE**. That is, reverse biasing the gate causes a depletion of carriers in the channel. This type of FET is often called a **TYPE A FET**. IGFET's can also be operated in the depletion mode and be classed as type A.

Since the gate is isolated from the rest of the device in the IGFET, the gate in an N-channel IGFET can also be made positive. This causes the generation or enhancement of carriers in the channel. Under these conditions the IGFET operates in the **ENHANCEMENT MODE**. A properly constructed IGFET can operate in both depletion and enhancement modes. This type of FET is called a **TYPE B FET**. With type B operation, the device can handle both positive and negative input signals, thus allowing a greater input signal swing. Some JFET's can also be operated in the depletion mode and allow signal swings into the enhancement modes. Some IGFET's can only operate in the enhancement mode. This type of FET is referred to as a **TYPE C FET**. Type C IGFET's are commonly used in switching circuits, while type A and B IGFET's are commonly used as amplifiers, oscillators, etc.

FET AMPLIFIERS

Both junction transistors and vacuum tubes may be connected in three different amplifier configurations. For junction transistors the configurations are common emitter, common base and common collector (emitter follower). For vacuum tubes the configurations are common cathode, common grid (grounded grid) and common plate (cathode follower). Three similar amplifier configurations are possible with field effect transistors. Each configuration has its advantages and disadvantages.

Common Source

The **COMMON SOURCE AMPLIFIER** configuration shown in Figure 11 is by far the most widely used FET amplifier configuration. The input signal is applied between gate and ground. The output signal appears between drain and ground. The source is grounded and is the common element between input and output.

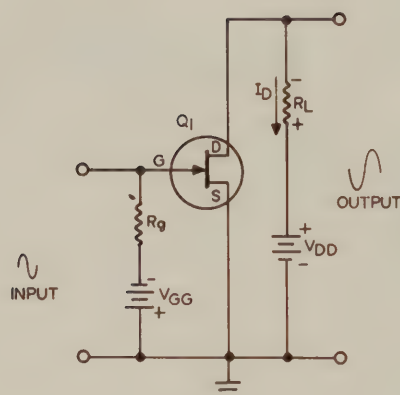


Figure
11

The gate bias supply V_{GG} makes the gate negative with respect to the source for proper operation (N-channel FET). The drain supply V_{DD} supplies the drain current. Gate resistor R_g develops the input signal and R_L , the load resistor, develops the output signal. This circuit is similar to the common emitter transistor amplifier and common cathode vacuum tube amplifier.

To examine circuit operation, suppose V_{GG} is equal to one half the pinch-off (cutoff) voltage V_P . This establishes a steady state value of drain current. Drain current flowing through R_L develops a voltage drop E_{R_L} of the polarity shown. The drain-source voltage and likewise the output voltage is the difference between V_{DD} and E_{R_L} .

Now suppose an input signal is applied to the circuit of Figure 11. As the input signal attempts to drive the gate positive, it opposes the negative bias set up by V_{GG} , and drain current increases. The voltage drop across R_L increases, thus making the drain-source voltage and the output voltage less positive. The opposite action occurs when the input signal drives the gate more negative. The negative input signal aids the negative bias set up by V_{GG} , and drain current is decreased, thus making the drain-source voltage and the output voltage more positive.

Summarizing the action, a positive input swing to an N-channel FET causes a decrease in output voltage (negative swing), and a negative input signal swing causes an increase in output voltage (positive swing). Thus, the common source amplifier produces a 180° phase shift between input and output.

In addition, the circuit has a very high input impedance (many megohms) since the gate PN junction is reverse biased. IGFET's have a still higher input impedance since the gate is like one plate of a capacitor, with the device itself being the other plate. The common source amplifier also provides a high voltage gain.

The circuit shown in Figure 11 employs an N-channel FET. The circuit is the same for a P-channel FET with the supply voltage polarities reversed, as shown in Figure 12. The 180° phase shift between input and output signals likewise occurs in the P-channel circuit.

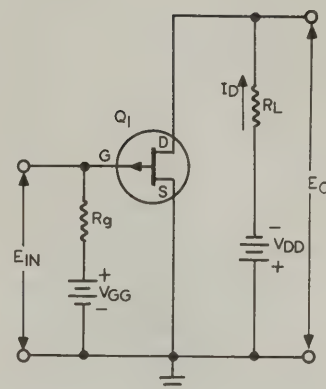


Figure 12

Common Drain

The **COMMON DRAIN AMPLIFIER** configuration is shown in Figure 13. As in the common source amplifier, the input signal is applied between the gate and ground. The output signal is developed across a load resistor connected in the source circuit. The drain is connected directly to the drain supply V_{DD} . From a signal standpoint, the positive terminal of V_{DD} is at ground potential. Thus, the drain terminal is common to both the input and output circuits.

Gate circuit bias is established by V_{GG} as in the common source amplifier. This circuit is similar to the common collector (emitter follower) transistor amplifier and the common plate (cathode follower) vacuum tube amplifier.

To examine circuit operation, again assume that V_{GG} sets the bias at about one half the pinch-off (cutoff) voltage. This establishes a steady state value of drain current which in turn establishes a steady state value of output voltage.

Now suppose the input signal goes positive. The positive input signal opposes the bias established by V_{GG} . This increases drain current, making the voltage drop across R_L and thus the output voltage more positive. The opposite action occurs when the input signal goes negative. The negative input signal aids the bias established by V_{GG} , thus reducing drain current. In turn, the voltage drop across R_L and thus the output voltage becomes less positive.

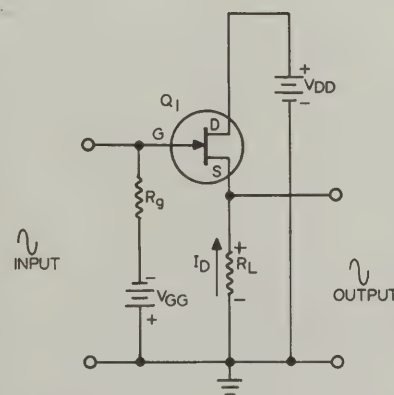


Figure 13

Thus, in the common drain amplifier the input and output signals are in phase. Since the input and output are in phase, the source "follows" the input. Thus, the circuit is often called a **SOURCE FOLLOWER**. A positive input signal swing causes the output to go more positive (positive swing), and a negative input signal swing causes the output to go less positive (negative swing).

The common drain or source follower amplifier has a very high input impedance and a low output impedance. The low output impedance makes the circuit useful for coupling a signal from a high impedance source to a low impedance conventional transistor circuit. In other words, the circuit makes an excellent impedance matching device. The voltage gain of the circuit is always less than 1, as is the case with its conventional transistor and vacuum tube counterparts.

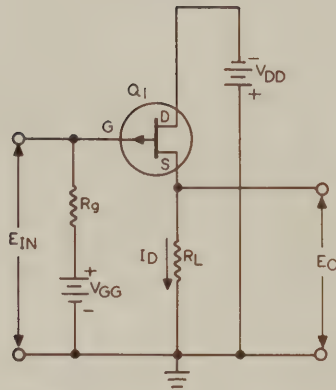


Figure 14

The circuit shown in Figure 13 employs an N-channel FET. A P-channel FET would require reversal of the supply polarities as shown in Figure 14.



This meter provides direct readout in testing either conventional bipolar or field effect transistors.

Courtesy Sencore, Inc.

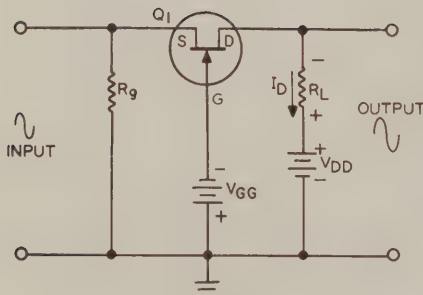


Figure 15

Common Gate

Figure 15 shows a **COMMON GATE AMPLIFIER** configuration. The input signal is applied between source and ground and the output is taken between drain and ground, with the gate element grounded. Thus, the gate is the common element between input and output circuits. As in the previous circuits, V_{GG} establishes the operating point. Drain current in R_L develops the output signal.

To examine circuit operation, suppose the input signal goes positive. Making the source positive is the same as making the gate negative. This action aids the bias established by V_{GG} , thus reducing drain current. In turn, the voltage drop across R_L is decreased. Since the output voltage is the difference between E_{R_L} and V_{DD} , the output voltage becomes more positive.

The opposite action occurs when the input signal goes negative. The negative input signal applied to the source is the same as a positive signal applied to the gate. This opposes the bias established by V_{GG} , thus increasing drain current. The voltage drop across R_L is increased, causing a decrease in the output voltage.

Thus, in the common gate amplifier the input and output signals are in phase. A positive input signal swing causes the output signal to become more positive (positive swing), while a negative input signal swing causes the output signal to become less positive (negative swing). The common gate amplifier has low circuit capacitance, which makes it useful in high frequency amplifier circuits. This is also true of its junction transistor and vacuum tube counterparts, the common base amplifier and common (grounded) grid amplifier.

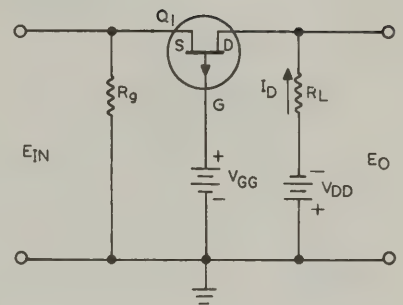


Figure
16

By reversing supply voltage and signal polarities, the circuit of Figure 15 can be set up for a P-channel FET as shown in Figure 16.

SELF BIAS

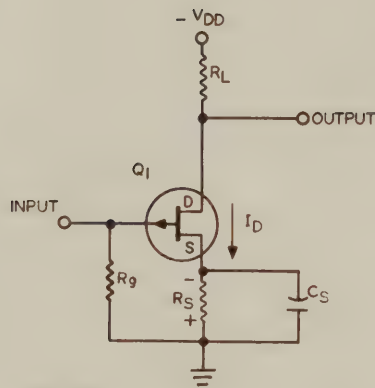


Figure 17

The use of a separate, fixed gate bias supply is expensive and often not desirable from a stability standpoint. FET characteristics vary from device to device, and a change in devices would almost always require a readjustment of the fixed supply. A **SELF BIAS** arrangement is more tolerant of device variations, and is also simpler. The self bias technique used on FET's is identical to the cathode or self bias used with vacuum tube amplifiers.

Figure 17 shows how self bias can be obtained in a P-channel common source amplifier circuit. Drain current flowing through R_s develops a voltage drop of the polarity shown. The gate is connected to the positive end of this resistor through R_g , while the source is connected to the negative end of R_s . Thus, the gate is made positive with respect to the source, providing the desired reverse bias for a P-channel FET. As far as bias is concerned, there is a negligible voltage drop across R_g since the gate is reverse biased and only a very small reverse current flows in the gate circuit.

Bypass capacitor C_s is connected across the bias resistor to prevent any signal variations from occurring across the resistor. Thus, C_s places the source at ground potential for signal frequencies.

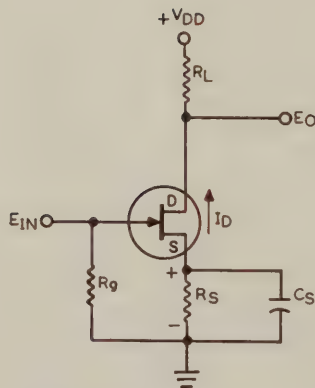


Figure 18

MULTI-STAGE AMPLIFIERS

By employing conventional coupling techniques (RC, transformer or impedance coupling), FET's can be used in multi-stage amplifiers. The high input impedance of the FET greatly eliminates the loading effect of the second stage on the first, which is a problem in multi-stage amplifiers that use conventional junction transistors. In many cases, FET's are used in the input stages, with conventional junction transistors forming the remainder of the amplifiers.

Figure 19 shows a multi-stage amplifier employing RC coupling with both FET's and a conventional junction transistor. Q_1 operates as a common source amplifier. R_3 and C_1 provide self bias for this stage. The signal at the drain of Q_1 is coupled to the gate of Q_2 through C_2 . The charge and discharge of C_2 develops the input signal for Q_2 across R_4 . Q_2 operates as a common drain or source follower amplifier. R_5 develops self bias for Q_2 in addition to developing the output signal.

The Q_2 source follower has a very low output impedance and thus provides an impedance match to the input circuit of Q_3 , a conventional NPN junction transistor. C_3 couples the signal to the base of Q_3 . Bias is established for Q_3 by R_6 and R_7 . The output signal is taken from the collector of Q_3 .

The following Practice Exercise questions cover the subjects which you have just studied. They are:

INSULATED GATE FET

OPERATING MODES

FET AMPLIFIERS

COMMON SOURCE

COMMON DRAIN

COMMON GATE

25. How is an insulated gate FET different from a junction FET?
26. Another name for an insulated gate FET is the (a) MOS, (b) JFET.
27. Draw the symbols for N-channel and P-channel IGFET's.
28. JFET's normally operate in the (a) depletion mode, (b) enhancement mode.
29. A depletion mode FET is said to be a type _____ FET.
30. What is a type B FET?
31. A JFET can operate type C. True or False?
32. In the circuit of Figure 11, the drain is the common element. True or False?
33. What phase shift is produced between the input and output signals in a common source amplifier?
34. Draw a common source amplifier using a P-channel FET.

35. The circuit shown in Figure 13 is a (a) common drain amplifier, (b) common source amplifier, (c) common gate amplifier.
36. A 180° phase shift is produced between input and output signals in a source follower amplifier. True or False?
37. Draw a source follower amplifier using a P-channel FET.
38. The input impedance of a source follower is _____ and the output impedance is _____.
39. A _____ degree phase shift occurs in a common gate amplifier.
40. Draw a common gate amplifier using a P-channel FET.

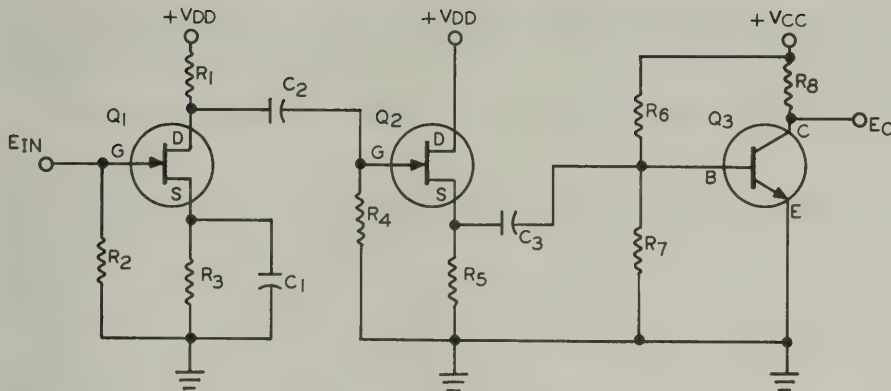


Figure 19

DIRECT-COUPLED AMPLIFIERS

Early types of multi-staged transistor amplifiers used conventional coupling techniques, as just discussed, because direct coupling was not practical. Later developments in field effect and planar transistors have made the circuit simplification of direct coupling a practical reality.

Multi-stage amplifiers can be formed by coupling the collector or drain directly to the input of the following stage. DC voltage and current sources must, of course, provide the proper bias levels required for each stage. With each additional stage, a conventional dc amplifier becomes increasingly difficult to design and adjust. It has been found convenient, however, to

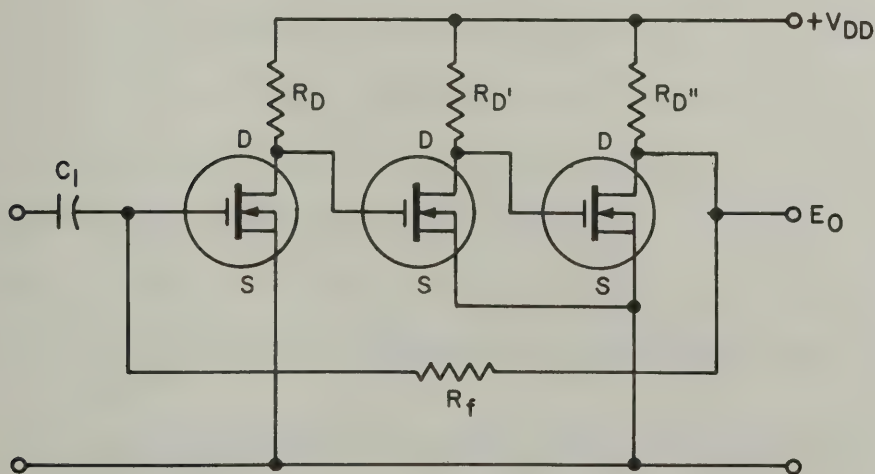


Figure 20

direct-couple transistor stages in pairs; using RC coupling or other dc isolation methods only between pairs. These pairs are also called **COMPOUNDS**, and refer to the various types of direct-coupled pairs of amplifiers. A pair or compound is sometimes made a triplet by adding a third stage.

One example which illustrates the design convenience of direct-coupled FET amplifiers uses the P-channel, enhancement-mode IGFET. The P-channel, enhancement-mode MOSFET requires a negative drain voltage and a negative gate bias that allows the use of very simple bias techniques. Figure 20 shows a circuit in which an N-channel MOSFET triplet is operated with the gate and drain at the same potential, making it possible to direct-couple a series of stages. The amplifier uses three identical stages with a feedback resistor connecting the output stage drain back to the input gate. The fact that the drain resistors are identical makes the operating points (Q-points) the same, and independent of the supply voltage.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

SELF BIAS

MULTI-STAGE AMPLIFIERS

DIRECT-COUPLED AMPLIFIERS

R-F AMPLIFIERS

41. In Figure 17, the self bias is developed across (a) R_g , (b) R_s , (c) R_L .
42. What is the purpose of C_s in Figure 17?
43. Draw an N-channel common source amplifier with self bias.
44. Q_1 in Figure 19 operates as a (a) common source amplifier, (b) common gate amplifier, (c) common drain amplifier.
45. Q_2 in Figure 19 operates as a (a) common source amplifier, (b) common gate amplifier, (c) common drain amplifier.
46. Q_2 in Figure 19 is a P-channel FET. True or False?
47. What coupling method is employed in Figure 19?
48. What bias method is used on Q_1 in Figure 19?
49. Q_1 in Figure 19 has a (a) high input impedance, (b) low input impedance.
50. Q_2 in Figure 19 has a (a) high output impedance, (b) low output impedance.
51. The circuit of Figure 21 employs an insulated gate field effect transistor. True or False?
52. The low input impedance of an FET results in a low circuit Q, thus providing a wide bandwidth. True or False?

R-F AMPLIFIERS

Field effect transistors can be used in r-f amplifier circuits. They are especially useful in the front end (r-f amplifier and mixer stages) of radio receivers. FET's are not as subject to overloading by strong signals as are conventional junction transistors and vacuum tubes. This is particularly true of the MOS type of FET.

Figure 21 shows a simple N-channel MOS (IGFET) r-f amplifier circuit. The input signal is coupled to the Q_1 gate through L_1 and L_2 . The high input impedance of the FET does not load down the tuned circuit, and thus permits high circuit Q and sharp selectivity. Self bias for Q_1 is provided by R_1 and C_2 . The output signal is coupled to the load through L_3 and L_4 .

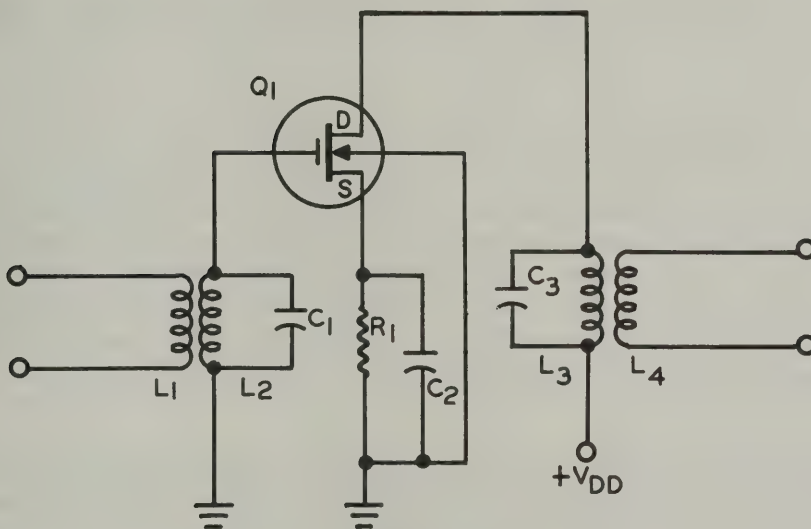


Figure 21

The amplifier circuits previously presented show the basic uses for field effect transistors. FET's find use in all types of electronic circuits and applications, from amplifiers and oscillators to switching circuits. Their prime use is, however, in circuits requiring very high input impedances.

SUMMARY

The field effect transistor, although it is called a transistor, operates in a fashion much different from the conventional junction transistor. In the junction FET, the reverse bias applied to the gate sets up depletion regions in the channel. These depletion regions control the flow of current from source to drain.

The three elements of an FET (source, gate and drain) are analogous to the three elements of a transistor (emitter, base and collector) and to the three basic vacuum tube elements (cathode, control grid and plate). In all three devices, current through the device is controlled by applying a voltage or current to the control element.

FET's can be constructed with either N-type or P-type channel material. Accordingly, the devices are referred to as N-channel or P-channel FET's.

The conventional FET is made with an actual PN junction at the gate element. This PN junction is reverse biased for normal operation, giving the FET a very high input impedance. The FET can also be constructed with the gate element actually insulated from the rest of the device by a layer of silicon oxide or other insulating material. This type of FET is referred to as an insulated gate FET (IGFET) or metal-oxide-semiconductor (MOS). The conventional FET, since it employs a PN junction, is often referred to as a junction FET (JFET).

JFET's are normally operated with the gate reverse biased. This type of operation is referred to as depletion mode. Some JFET's, but mostly IGFET's can be operated with the gate forward biased also. This type of FET then operates in the depletion and enhancement modes. An IGFET can also be operated strictly in the enhancement mode.

FET's can be connected in three basic amplifier configurations: common source, common drain (or source follower) and common gate. Each type has its advantages and disadvantages. The common source has the highest voltage gain. The common drain or source follower is an impedance matching circuit with high input impedance and low output impedance.

IMPORTANT DEFINITIONS

AMPLIFICATION FACTOR (μ) — The theoretical maximum voltage amplification of an FET.

CHANNEL — The controlled region through which current flows in a field effect transistor.

CHANNEL PINCH-OFF — A condition in a field effect transistor where the depletion regions almost meet and permit no further increase in drain current.

COMMON DRAIN AMPLIFIER — An FET amplifier configuration in which the drain is common to both input and output circuits. Also called a SOURCE FOLLOWER.

COMMON GATE AMPLIFIER — An FET amplifier configuration in which the gate is common to both input and output circuits.

COMMON SOURCE AMPLIFIER — An FET amplifier configuration in which the source is common to both input and output circuits.

DEPLETION MODE — A type of FET operation in which carriers are removed or depleted from the channel for current control. This is the normal operation for junction FET's.

DRAIN — The element in a field effect transistor which collects current, analogous to the collector of a junction transistor or plate of a vacuum tube.

ENHANCEMENT MODE — A type of insulated gate FET operation where carriers are added to the channel for current control.

FIELD EFFECT TRANSISTOR (FET) — A semiconductor device whose output current is controlled by an electric field within the device.

FORWARD TRANSFER ADMITTANCE (y_{fs}) — The ability of the gate-source voltage to control drain current in an FET.

GATE — The control element in a field effect transistor, analogous to the base of a junction transistor or control grid of a vacuum tube.

INSULATED GATE FET (IGFET) — A type of FET in which the gate is insulated from the channel by a thin layer of silicon oxide, nitride, or other insulating material.

JUNCTION FET (JFET) — The conventional type of FET in which the gate forms a PN junction with the channel.

IMPORTANT DEFINITIONS (Continued)

METAL-OXIDE-SEMICONDUCTOR (MOS) — Another name for the insulated gate FET.

N-CHANNEL FET — A field effect transistor in which the conducting material is N-type semiconductor material.

OUTPUT ADMITTANCE (y_{os}) — The reciprocal of the output impedance of an FET.

P-CHANNEL FET — A field effect transistor in which the conducting material is P-type semiconductor material.

SELF BIAS — The use of a source resistor in an FET amplifier circuit to provide bias for the gate circuit.

SOURCE — The source of current in a field transistor analogous to the emitter of a junction transistor or cathode of a vacuum tube.

SOURCE FOLLOWER — See COMMON DRAIN AMPLIFIER.

TYPE A FET — A field effect transistor which operates in the depletion mode only.

TYPE B FET — A field effect transistor which operates in both the depletion and enhancement modes.

TYPE C FET — A field effect transistor which operates in the enhancement mode only.

ESSENTIAL SYMBOLS AND EQUATIONS

V_{DS} — drain-source voltage	(volts)
V_{GS} — gate-source voltage	(volts)
I_D — drain current	(amperes)
I_{DSS} — drain saturation current	(amperes)
V_P — gate-source pinch-off voltage	(volts)
I_{GSS} — gate-source cutoff current	(amperes)
V_{DD} — drain supply voltage	(volts)
V_{GG} — gate supply voltage	(volts)
y_{fs} — forward transfer admittance	(mhos, siemens)
y_{os} — output admittance	(mhos, siemens)
Z_o — output impedance	(ohms)
μ — amplification factor	

$$y_{fs} = \frac{\Delta I_D}{\Delta V_{GS}} \text{ with } V_{DS} \text{ held constant} \quad (1)$$

$$y_{os} = \frac{\Delta I_D}{\Delta V_{DS}} \text{ with } V_{GS} \text{ held constant} \quad (2)$$

$$\mu = \frac{\Delta V_{DS}}{\Delta V_{GS}} \text{ with } I_D \text{ held constant} \quad (3)$$

$$\mu = \frac{y_{fs}}{y_{os}} \quad (4)$$

PRACTICE EXERCISE SOLUTIONS

1. (a) high input impedance.
2. False — The FET's operation is closer to that of a vacuum tube than a conventional junction transistor.
3. Source, drain and gate.
4. False — PN junctions are only formed at the gate connections.
5. (b) reverse biased.
6. The channel is that area between the gate PN junctions that is not a part of the depletion regions.
7. Channel pinch-off is the point where the depletion regions almost meet in the channel, causing current I_D to reach its saturated value.
8. (a) lower value of drain current.
9. (b) decrease.
10. The symbol for an N-channel FET is shown in Figure 4A, and the symbol for a P-channel FET is shown in Figure 4B.
11. base
12. (a) single-ended.
13.
 - a. V_{GS} — gate-source voltage
 - b. V_{DS} — drain-source voltage
 - c. I_{DSS} — drain current at saturation
 - d. BV_{DGS} — drain-gate breakdown voltage
 - e. I_D — drain current
 - f. I_{GSS} — gate-source cutoff current
14. saturation
15. Pinch-off voltage (V_P) is the gate-source voltage V_{GS} necessary to reduce drain current I_D to a low, cutoff value.
16. cutoff voltage
17. y_{fs} is the forward transfer admittance and is found by dividing a change in I_D by the change in V_{GS} that produced the change in I_D . The measurement is made with V_{DS} held constant.

PRACTICE EXERCISE SOLUTIONS (Continued)

18. (c) microsiemens or micromho.
19. transconductance (g_m)
20. y_{os} is the output admittance and is found by dividing a change in I_D by the change in V_{DS} that produced the change in I_D . The measurement is made with V_{GS} held constant.
21. 50 k Ω — The output impedance is found by taking the reciprocal of y_{os} :

$$Z_o = \frac{1}{y_{os}} = \frac{1}{.02 \times 10^{-3} \text{ mhos}} = 50 \times 10^3 \text{ ohms} = 50 \text{ k}\Omega$$

22. dynamic plate resistance (r_p)
23. μ is the amplification factor and is found by dividing a change in V_{DS} by the change in V_{GS} that produced the change in V_{DS} . The measurement is made with I_D held constant.

24. $\mu = \frac{y_{fs}}{y_{os}}$

25. The gate electrode does not form a PN junction and is insulated from the rest of the device.
26. (a) MOS.
27. The symbol for an N-channel IGFET is shown in Figure 10A, and the symbol for a P-channel IGFET is shown in Figure 10B.
28. (a) depletion mode.
29. A
30. A type B FET is one that is designed to operate in the depletion and enhancement modes.
31. False — The gate is normally reverse biased and only depletion mode operation is usually obtainable. Some JFET's allow signal swings into the enhancement region. However, they normally are not forward biased.
32. False — The circuit of Figure 11 is a common source amplifier.
33. 180°

PRACTICE EXERCISE SOLUTIONS (Continued)

- 34. A common source amplifier using a P-channel FET is shown in Figure 12. Notice that the bias polarities and the direction of drain current are opposite those in the N-channel circuit of Figure 11.
- 35. (a) common drain amplifier.
- 36. False — The input and output signals are in phase in a source follower amplifier.
- 37. A source follower using a P-channel FET is shown in Figure 14. The circuit arrangement is similar to that shown for the N-channel FET, except for the reversal of bias polarities and thus, the reversal in direction of drain current.
- 38. high, low.
- 39. zero — The input and output signals are in phase.
- 40. A common gate amplifier using a P-channel FET is shown in Figure 16. The arrangement is similar to that of the N-channel circuit except for the reversal of bias polarities, and thus the reversal of drain current.
- 41. (b) R_s .
- 42. Bypass capacitor C_s prevents any signal frequencies from appearing across R_s , thus placing the source at ground as far as signal frequencies are concerned.
- 43. An N-channel common source amplifier using self bias is shown in Figure 18. Drain current flowing upward through source resistor R_s produces a voltage drop which serves as bias for the FET.
- 44. (a) common source amplifier.
- 45. (c) common drain amplifier.
- 46. False — Q_1 and Q_2 in Figure 19 are N-channel FET's.
- 47. RC coupling
- 48. self bias
- 49. (a) high input impedance.
- 50. (b) low output impedance.

PRACTICE EXERCISE SOLUTIONS (Continued)

51. True

52. False — The FET has a high input impedance, thus providing minimum loading and high Q.



1100A

1. A - IN A JUNCTION FIELD EFFECT TRANSISTOR, A PN JUNCTION IS FORMED AT THE -- GATE.

The gate is formed by diffusing opposite-type impurities into a block of P or N-type silicon, and thus a junction is formed at the boundary between the two dissimilar materials. The connections to the source and drain are ohmic (purely resistive) contacts.

2. A - FOR AN N-CHANNEL JFET -- THE GATE IS MADE NEGATIVE.

Since the gate is a P-diffused region in an N-channel JFET, and since the gate represents a reverse-biased diode, the polarity required is a negative V_{GS} .

3. A - THE CONDITION IN AN FET WHERE THE DEPLETION REGIONS ALMOST MEET IS CAUSED BY -- THE PINCH-OFF VOLTAGE.

A reverse-biased diode, such as the gate in an FET, produces depletion regions in which current carriers are not available. Since this depletion occurs in the drain-current channel, the back-biasing can be increased until there are no current carriers available in the channel, and drain current becomes zero. This effect is analogous to the cutoff condition in a vacuum tube.

4. B - FIGURE 4B SHOWS -- A P-CHANNEL JFET.

The drain of a P-channel JFET must be connected to a negative supply, and the arrow on a JFET points AWAY from a negative drain.

5. A - THE VALUE OF DRAIN CURRENT AT SATURATION IS CALLED -- I_{DSS} .

The drain saturation current is the leveling-off point on an FET characteristic curve in which further increases in the drain-source voltages produce only small increases in drain current. This is also the channel pinch-off condition for a given V_{GS} .

6. B - THE FORWARD TRANSFER ADMITTANCE OF AN FET IS CALLED -- y_{fs} .

y_{fs} is the measure of the change in drain current caused by a given change in gate voltage.

7. C - THE INSULATED GATE FET EMPLOYS -- NO PN JUNCTION AT THE GATE.

The IGFET uses an insulating layer between the metalization layer and the channel material, which forms a small, high-quality capacitor at the gate.

8. D - JFET's NORMALLY OPERATE -- IN THE DEPLETION MODE.

JFET's require the back-biased gate diode which causes depletion regions in the channel. The depletion of current carriers in the channel controls the drain current, thus the JFET operates in the depletion mode.

9. C - AN FET EXHIBITS A CHANGE IN I_D OF .75 mA WITH A 5 VOLT CHANGE IN V_{DS} . THE OUTPUT ADMITTANCE IS -- 150 MICROMHOS.

The output admittance y_{os} is found by dividing the change in drain current by the change in V_{DS} .

$$\begin{aligned} y_{os} &= \frac{\Delta I_D}{\Delta V_{DS}} \\ &= \frac{.75 \times 10^{-3} \text{ amps}}{5 \text{ volts}} \\ &= .15 \times 10^{-3} \text{ mhos or } 150 \text{ } \mu\text{mhos} \end{aligned}$$

10. D - THE FET BIAS ARRANGEMENT SHOWN IN FIGURE 17 IS REFERRED TO AS -- SELF BIAS.

Self bias is analogous to cathode bias used in vacuum tube circuits. The RC network provides a voltage drop which makes the gate positive with respect to the source.

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1100A

Example: An animal commonly found in the jungles of South America is the
 A ☐ B ☒ C ☐ D ☐
 A. walrus. B. jaguar. C. polar bear. D. penguin.

1. A ☒ B ☐ C ☐ D ☐ In a junction field effect transistor, a PN junction is formed at the
 (A) gate. (B) source. (C) drain. (D) base.
2. A ☒ B ☐ C ☐ D ☐ For an N-channel JFET
 (A) the gate is made negative. (B) the drain is made negative. (C) the gate is connected to the source. (D) the gate is made positive.
3. A ☒ B ☐ C ☐ D ☐ The condition in an FET where the depletion regions almost meet is caused by
 (A) the pinch-off voltage. (B) conduction. (C) forward bias. (D) enhancement.
4. A ☐ B ☒ C ☐ D ☐ Figure 4B shows
 (A) an N-channel FET. (B) a P-channel FET. (C) an MOS. (D) an IGFET.
5. A ☒ B ☐ C ☐ D ☐ The value of drain current at saturation is called
 (A) I_{DSS} . (B) y_{fs} . (C) I_D . (D) I_{GSS} .
6. A ☐ B ☒ C ☐ D ☐ The forward transfer admittance of an FET is called
 (A) y_{os} . (B) y_{fs} . (C) μ . (D) I_{DSS} .
7. A ☐ B ☐ C ☒ D ☐ The insulated gate FET employs
 (A) a PN junction at the gate. (B) a PN junction at the source. (C) no PN junction at the gate. (D) a PN junction at the drain.
8. A ☐ B ☐ C ☐ D ☒ JFET's normally operate
 (A) in the enhancement mode. (B) in the depletion and enhancement modes. (C) with the gate forward biased. (D) in the depletion mode.
9. A ☐ B ☐ C ☒ D ☐ An FET exhibits a change in I_D of .75 mA with a 5 volt change in V_{DS} . The output admittance is
 (A) 666 micromhos. (B) 375 micromhos. (C) 150 micromhos. (D) 75 micromhos.
10. A ☐ B ☐ C ☐ D ☒ The FET bias arrangement shown in Figure 17 is referred to as
 (A) fixed bias. (B) cathode bias. (C) voltage divider bias. (D) self bias.

SEMICONDUCTOR CONTROL DEVICES

1103

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





Semiconductor control devices are used extensively in power supplies, regulators, and electronic switching equipment. The unique properties of these semiconductors have revolutionized the design of control instrumentation. Shown above is a high voltage power supply with its control circuitry.

Courtesy Sorensen Div. Raytheon Company

CONTENTS

Two-layer Control Devices	Page 4
Breakdown Diodes	Page 5
Four-layer Control Devices	Page 9
Gate Controls	Page 12
Unijunction Transistors	Page 14
Electronic Switching	Page 17
Basic PNP Switching Circuits	Page 18
Turn-on and Turn-off Methods	Page 19
SCR Switching Circuits	Page 21
SCR Control Circuits	Page 23
Controlled Rectifier Ratings	Page 27
Triac Control Circuits	Page 28

*Every man his own ancestor, and every man his own heir.
He devises his own future, and he inherits his own past.
—H. F. Hedge*

SEMICONDUCTOR CONTROL DEVICES

Industrial processing, scientific research, automatic computing and data handling, space technology, communications, and many other fields employ electrically operated instruments and machines. Each particular unit or system uses its electric power in the form of current at some specific voltage. Often, it is necessary that the voltages, currents, or both remain very near their specified values. Also, there is frequent need to change the supply current from ac to dc, or to change the voltage to a higher or lower value than that provided by the power lines.

These and other needs can be filled by means of various semiconductor devices that we refer to as control devices. Due to the small size and high efficiency of some of the semiconductor control devices, many new things are happening. Small control units are taking the place of older equipment that was so large it needed specially reinforced buildings for installation. Huge motors are being controlled with small semiconductor control devices. This lesson describes some commonly used devices of this kind. Their principles of operation are relatively simple. However, despite their simplicity, they are highly important because they make possible the proper operation of instruments and machines in practically every field of endeavor.

TWO-LAYER CONTROL DEVICES

The PN junction diode is one of the simplest of semiconductor devices employed in control applications. In each of these units, a single rectifying junction is formed at the junction of the P-type and an N-type semiconductor materials. A voltage applied across the terminals of such a diode is called a FORWARD BIAS when it makes the P-material or anode positive with respect to the N-material or cathode. Electrons then flow through the diode easily, and the current thus produced is called FORWARD CURRENT. An applied voltage of the opposite polarity is called REVERSE VOLTAGE or bias, and the current it produces is called REVERSE CURRENT.

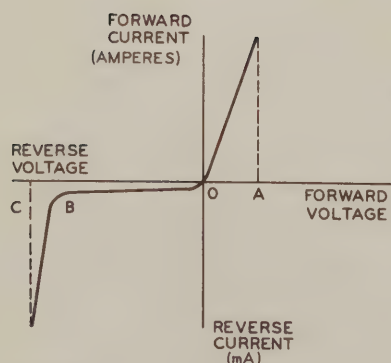


Figure 1

Figure 1 shows the forward and reverse characteristics of a PN diode. Notice that forward current is labeled in amperes and reverse current in milliamperes. For small current diodes the forward current would be in mA and the reverse current in μA . If the reverse current scale were the same as the forward current scale, the leakage current curve would lie on the horizontal axis of the graph. The part of the curve between O and A shows that forward current rises sharply to high values as forward voltage increases. This indicates that the diode has low forward resistance. For reverse voltages from O to B, the slope of the curve is very slight and the reverse current values in this

range are very low, indicating high back or reverse resistance. In most diode applications, applied voltages are limited so that operation is kept within the region between A and B. Here, the back resistance may be as great as a hundred million times the forward resistance. Desired diode conduction then consists of forward current, while reverse current usually is not desirable.

BREAKDOWN DIODES

At the lower left in Figure 1, the curve shows that increasing the reverse voltage from B to C causes a very large increase of reverse current. That is, when the reverse voltage is increased beyond B, the back resistance suddenly becomes very low. Thus, a diode's back resistance may be several megohms when the applied reverse voltage is between O and B, but suddenly falls to only a few ohms when the voltage becomes greater than B. In the range between B and C, a very large change of current can occur with only a very small change of voltage across the diode.

The rather sudden change of back resistance near point B is called BREAKDOWN. The region between B and C is called the BREAKDOWN REGION. The large current that can be produced here does not damage the diode unless the maximum power rating of the diode is exceeded.

The doping of the semiconductor material determines the reverse voltage at which breakdown occurs. In the manufacturing process, the resistivity can be controlled so as to select various desired breakdown voltages.

Diodes intended for operation in the breakdown region are often employed for control applications. These units are called BREAKDOWN DIODES. Other designations used are ZENER DIODE and AVALANCHE DIODE.

Technically, breakdown that occurs at a voltage of about 6 volts or less generally is considered to be due to an internal action known as tunneling or zener breakdown. Avalanche breakdown is the process responsible for breakdown at higher voltages. However, it is common practice to refer to these diodes as zener diodes, regardless of their breakdown voltage. Also, the current conducted in the breakdown region is called ZENER CURRENT (I_Z).

The control action provided by a zener diode is due to the important fact that, in the breakdown region, large current changes are associated with very small voltage changes. As a typical example, Figure 2 shows the reverse characteristics of a particular zener diode. This unit carries a rated **BREAK-**

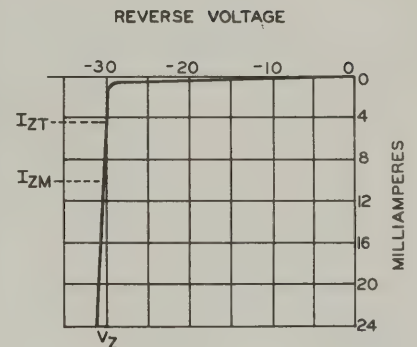
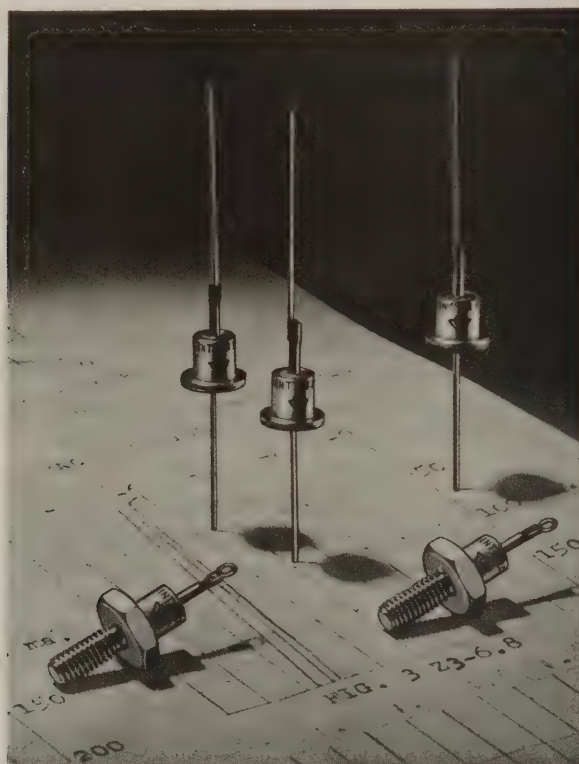


Figure
2

DOWN VOLTAGE of -30 volts. This value is the reverse voltage across the diode when the zener current is equal to some specified value, I_{ZT} , called the **TEST CURRENT**. The higher current indicated, I_{ZM} , is the maximum zener current that can be permitted without exceeding the power dissipating ability of the diode. In this example, I_{ZM} is equal to 10 mA.

Test current I_{ZT} is selected to approximate the zener current values most commonly used in circuit applications, and therefore is well below the value of I_{ZM} . In this example, I_{ZT} is 4.2 mA. Notice that the rated breakdown voltage of -30 volts is somewhat higher than the reverse bias value at which the breakdown action first begins to occur. This rated breakdown voltage is also referred to as the zener voltage, V_Z .

A common application of the zener or avalanche diode is voltage regulation. The voltage to be regulated may be the operating bias supplied to an element of a transistor or electron tube, or to some other circuit component.



A group of zener diodes. The three with pigtail leads are rated at 1.0 watt. The stud-mounted units have 3.5 watt ratings.

Courtesy International Rectifier Corp.

Figure 3 shows a voltage regulating circuit composed of a resistor and a zener diode. Notice that the symbol for the zener diode is different from that of a regular diode. It is identified by the diagonal lines at the ends of the cathode bar. This is one of the symbols used for zener diodes. In Figure 3, V_{IN} is the unregulated input voltage from the supply, while V_O is the regulated output voltage to be applied to the element or device.

Resistor R_1 and zener diode D_1 form a voltage divider across which V_{IN} divides into two parts, V_1 and V_Z . As explained, zener voltage V_Z is the voltage across the diode at a specified reverse current in the breakdown region. Thus, V_Z would be 30 volts if the diode represented in Figure 2 were employed in Figure 3.

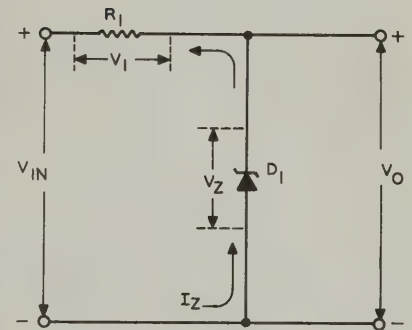


Figure 3

The anode of D_1 connects to the negative terminal of the voltage supply. Therefore, the anode is negative with respect to its cathode by the amount of the voltage V_Z . The resulting reverse or zener current I_Z consists of electron flow in the direction indicated by the arrows. As the output voltage is taken from across the diode terminals, V_O is equal to V_Z and therefore is less than V_{IN} by the amount of V_1 .

Briefly, the action of this regulator is such that a change of V_{IN} results in a like change of V_1 so that V_Z and V_O remain almost completely unchanged. For example, suppose V_{IN} increases. This tends to increase both V_1 and V_Z . However, an increase in V_Z of only a small fraction of a volt causes the resistance of the zener diode to decrease by a relatively large amount. This resistance is then a smaller fraction of the total resistance of R_1 and D_1 . As a result, output voltage V_O becomes a smaller fraction of V_{IN} , while V_1 becomes a larger fraction of V_{IN} . The increase in V_1 is almost equal to the increase in V_{IN} . Hence, the increase in V_O is very small.

A like action takes place when V_{IN} decreases. In doing so, it tends to decrease both V_1 and V_Z . But a very small decrease in V_Z makes the resistance of the diode increase by a relatively large amount. This changes the voltage division so that V_O becomes a larger fraction of V_{IN} , and V_1 becomes a smaller fraction of V_{IN} . Again, the change in V_1 is almost equal to the change in V_{IN} , and the decrease in V_O is very small. Thus, the regulator circuit action causes output voltage V_O to remain nearly constant when V_{IN} varies about its desired value.

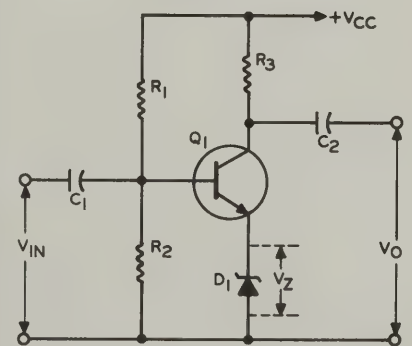


Figure 4

Figure 4 shows another voltage regulating application of the zener diode. Here, diode D_1 prevents changes in the emitter voltage of transistor Q_1 . The circuit, an amplifier, produces output V_O in the form of an amplified version of V_{IN} . Efficient amplifying action requires that the Q_1 emitter current change without causing changes of emitter voltage.

With the arrangement shown, the Q_1 emitter current is also the zener current of D_1 . As shown in Figure 2, the zener current can vary a relatively large amount with very little change in the voltage across the diode. This voltage, V_Z , is also the emitter voltage of Q_1 in Figure 4. Thus, the emitter voltage V_Z remains practically constant when the amplifier action causes changes in the emitter current.

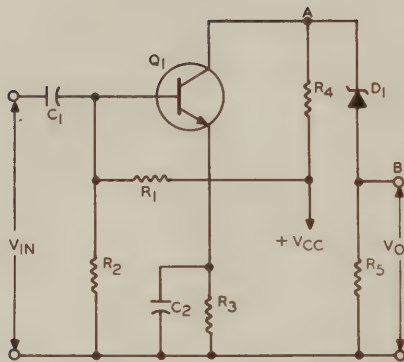


Figure 5

In Figure 5, zener diode D_1 couples the output signal from the collector of Q_1 to the input of a following stage. This arrangement permits direct coupling between the two stages even though the dc operating voltages at points A and B differ. To reverse bias the diode, D_1 is inserted so that its cathode connects to the more positive of the two points, which is point A in this case. A diode is selected which has a zener voltage equal to the difference of potential between A and B. This voltage produces a continuous zener current which consists of electron flow from the common connection up through R_5 and D_1 to point A.

Signal variations in the voltage at point A are applied to D_1 and R_5 in series. This causes variations in the zener current, and voltage changes across the diode and R_5 . However, the zener voltage across D_1 changes only very slightly, so that most of the variation is across R_5 . This is like the circuit of Figure 3, except that in Figure 5 the positions of the diode and resistor are reversed so that output V_O is the voltage across R_5 . Thus, any voltage change at A is coupled by the zener diode to B. As the change in voltage across D_1 is extremely small, the variation produced at B is practically equal to that at A.

Zener diodes are manufactured with different power handling abilities for various applications. Usually, the power rating is equal or approximately equal to the product of the rated zener voltage and the maximum allowable dc zener current. For example, a diode rated at $E_{ZT} = 4$ volts and $I_{ZM} = 250$ mA has a power rating of 1 watt. Typical power ratings include 250, 400, 500, and 750 milliwatts, and 1, 1.5, 3.5, 10, and 50 watts.

In use, the power that a diode is called upon to dissipate is generally much lower than its power rating. Each unit has a test current (I_{ZT}) rating, near which normal operation is recommended by the manufacturer. Thus, if the unit of the above example has an I_{ZT} rating of 50 mA, operation at this zener current level results in a power dissipation of only $4 \times .050 = 0.2$ watt, or one-fifth of its power rating.

Besides maximum power dissipation and zener voltage, zener diodes are rated as to their dynamic impedance, maximum allowable zener current, operating temperature range and, as mentioned, test current, which is the typical operating zener current. Other ratings sometimes given include the

saturation current (I_s), which is the leakage or reverse current at a reverse voltage considerably less than the zener breakdown voltage, and the maximum forward voltage (V_{FM}) that may be applied at a stated value of forward current. Typical values of I_s range from 5 to 1700 microamperes, while a typical rated V_{FM} is 1.5 volts at 200 mA.

The **DYNAMIC IMPEDANCE (Z_z)** or **DYNAMIC RESISTANCE (R_z)** of a zener diode is a measure of the diode's opposition to changes of zener current. The rated dynamic impedance, Z_z , is calculated from the Ohm's Law relation:

$$Z_z = V_{AC}/I_{AC}$$

where V_{AC} is the ac voltage required to superimpose an alternating current, I_{AC} , upon the rated dc test current. Usually, the rms value of I_{AC} is made equal to 10 per cent of test current I_z .

The dynamic impedance or resistance is one of the most important factors to be considered when selecting a zener diode for use in a voltage regulator application. All other things being equal, the lower the dynamic impedance the better the regulating action. For various zener diodes, rated Z_z values range from about 1 ohm to 2500 ohms or more. The Z_z of a given diode varies with the dc zener current. At large I_z values, Z_z is small, while for small values of I_z , the dynamic impedance is large.

Typical rated operating temperature ranges are from -65°C to 175°C junction temperature. The rated maximum power dissipation is generally specified with reference to an ambient temperature of 25°C . Often this rating holds for temperatures as high as 35°C to 70°C also. However, beginning at some temperature between 25°C and 70°C , for various units, the maximum power rating falls with an increase in temperature. Thus, with a given zener voltage, the maximum allowable zener current decreases with a rise in ambient temperature.

FOUR-LAYER CONTROL DEVICES

Just as a P-type semiconductor and an N-type semiconductor form a two-layer PN device, or diode, other semiconductor devices are built up of three or more alternate layers of P and N materials. Most transistors are three-layer units of either PNP or NPN structure, and therefore contain two PN junctions.

Four layers provide a PNPN device, which has three PN junctions. A large number of applications have been found for various four-layer devices. These units have characteristics which make them superior to the two-layer and three-layer devices for certain control actions.

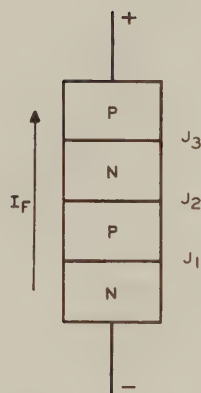


Figure 6

Figure 6 shows the basic PNPN arrangement. The three junctions are indicated as J_1 , J_2 and J_3 . This device is said to be forward biased when an applied voltage makes the outer P-region positive and the outer N-region negative, as indicated. Each junction has a certain resistance. In series, these resistances act as a voltage divider so that part of the applied voltage is across each junction. Each of the upper three regions is more positive than the one just below it. Therefore, the voltages across junctions J_1 and J_3 forward bias these junctions, while that across J_2 is a reverse bias.

Since the three junctions are in series, any current carried by the entire PNPN unit must pass through all of the junctions. Forward biased junctions J_1 and J_3 present very little resistance to forward current, which is electron flow in the direction of arrow I_F . However, in passing through reverse biased junction J_2 , this flow is reverse current. Therefore, the total current is limited to the relatively small value permitted by this junction.

Figure 7 shows the voltage-current characteristics of the PNPN device. In the upper right quadrant, the horizontal part of the solid-line curve shows that forward current I_F remains at a low value as forward bias V_F increases from zero. This is due to the blocking action of the middle junction. This region of the curve is called the **FORWARD BLOCKING REGION**.

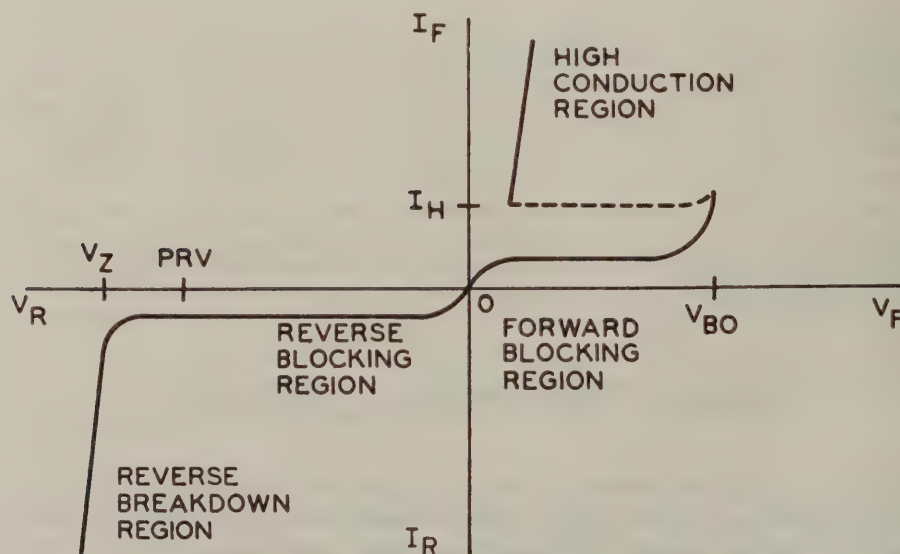


Figure 7

As V_F reaches high values, the current begins to increase; slowly at first, and then faster. This causes more and more current carriers to enter the J_2 barrier region, reducing its resistance. Finally, at voltage V_{BO} , this junction resistance suddenly becomes very low. This sudden change is called **FORWARD BREAKDOWN**. After it has occurred, relatively low forward voltages can produce very large forward currents. This is indicated by the nearly vertical part of the curve, labeled **HIGH CONDUCTION REGION**.

The voltage at which forward breakdown occurs is called the **BREAKOVER VOLTAGE (V_{BO})**. After breakdown has taken place, the total resistance of the PNP device is so low that I_F is limited primarily by the external circuit. So long as this circuit permits the current to be equal to or greater than a certain value called the holding current (I_H), the device continues to conduct in its forward breakdown region. Whenever I_F is reduced to less than I_H , the device "shuts off", quickly returning to its forward blocking region.

Referring to Figure 6, the PNP device may be reverse biased by making the upper P-region negative and the lower N-region positive. This causes junctions J_1 and J_3 to be reverse biased, and junction J_2 to be forward biased. In Figure 7, the reverse characteristics are shown in the lower left quadrant. Reverse current I_R remains low until the reverse bias reaches the breakdown value V_Z , then increases sharply with small increases of voltage. This is the same as the reverse characteristic of the PN junction diode. As indicated, the rated peak reverse voltage, PRV, is somewhat less than the reverse breakdown voltage V_Z .

There is an important difference between the forward and reverse characteristics of the PNP device. If a reverse bias much greater than the PRV rating is applied, the resulting large reverse current may destroy the device. However, a forward bias that exceeds the breakover voltage V_{BO} does not destroy the unit. It merely causes the device to switch over from the blocking to the high conduction region. In various applications, this device is commonly made to switch back and forth between these two regions or operating states. When operating in either blocking region, the device is said to be in its "off" state because the current it permits is very small. It is said to be in its "on" state when it is operating in the region of high conduction.

GATE CONTROLS

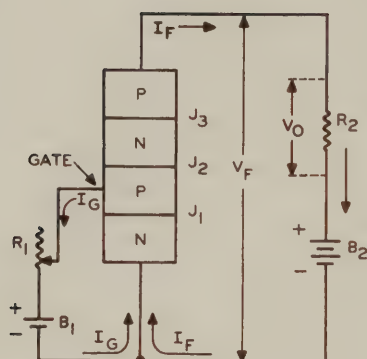


Figure 8

With the PNPN units of Figure 6, the only way to make it switch from the blocking or off state to the high conduction or on state is to increase the forward bias to the breakover voltage. Two-terminal units operated in this way are called four-layer diodes. Addition of a third terminal increases the range of applications of the PNPN device. As shown in Figure 8, the third terminal is brought out from the internal P-region. Usually, this terminal is called the **GATE**.

In the circuit of Figure 8, battery B_1 produces electron flow in the direction of arrow I_G . This current, called the gate current, flows from the negative terminal of B_1 to the lower N-region, through junction J_1 to the internal P-region, and from the gate terminal, through R_1 to the positive terminal of B_1 . The gate current can be varied by adjusting series resistor R_1 .

The voltage of battery B_2 is divided into two parts, V_F and V_O . Voltage V_F makes the upper P-region positive and the lower N-region negative, and thus forward biases the entire PNPN unit as in Figure 6.

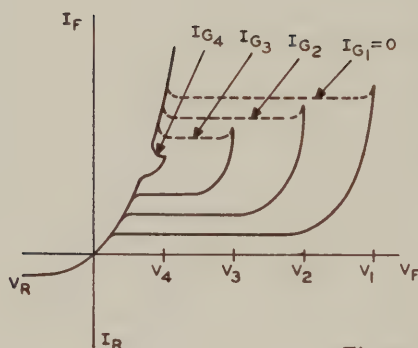


Figure 9

The forward breakover voltage depends upon the gate current. Figure 9 shows this relationship. When the gate current is zero (I_{G1}), the breakover voltage is highest, indicated as V_1 . Gate current I_{G2} of some low value causes the breakover voltage to be at V_2 . A larger gate current, I_{G3} , reduces the breakover point to V_3 . Finally, a relatively high gate current, I_{G4} , causes the breakover action to occur when V_F is equal to only V_4 .

Because of this relationship between gate current I_G and breakover voltage V_{BO} , the gate current circuit provides another means by which the PNPN device can be switched from the off state to the on state. This switching action will be explained a little later.

In Figure 8, voltage V_O is at all times equal to the B_2 battery voltage minus V_F . When the PNPN device is in its off state, its resistance is high. Then V_F is relatively large and V_O relatively small. When the device switches into its on state, the voltage across it suddenly falls to a relatively small value, and V_O becomes relatively large.

For example, referring to Figure 9, the change from the off to the on state may cause V_F to drop from V_2 to a value in the neighborhood of V_4 . If the battery voltage is equal to V_1 , then when the device is in the off state, V_O is equal to $V_1 - V_2$, a small value. When the device is in the on state, V_O is approximately equal to $V_1 - V_4$, a relatively large value.

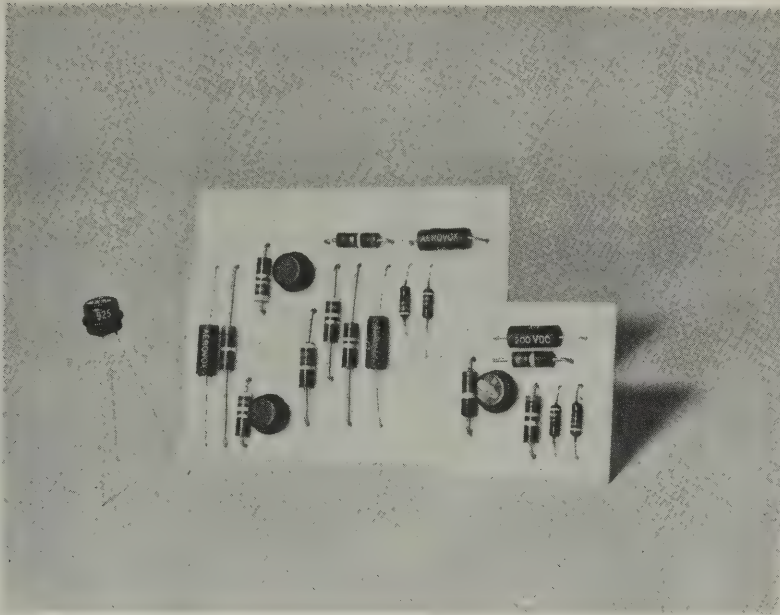
The following Practice Exercise questions cover the subjects which you have just studied. They are:

TWO-LAYER CONTROL DEVICES

BREAKDOWN DIODES

FOUR-LAYER CONTROL DEVICES

1. In Figure 1, the part of the curve between A and B shows that the (a) reverse current rises sharply, (b) reverse resistance is greater than the forward resistance, (c) reverse current is produced by forward voltage.
2. A zener diode is used in the (a) forward conduction region, (b) regions between zero bias and breakdown, (c) breakdown region.
3. In the zener diode breakdown region, large current changes can be produced by very small voltage changes. True or False?
4. A diode's zener voltage is (a) the rated breakdown voltage, (b) the largest permissible forward voltage.
5. The circuit of Figure 3 provides voltage regulating action because a small change in applied zener voltage causes a relatively large change in the resistance of (a) resistor R_1 , (b) diode D_1 , (c) both R_1 and D_1 .
6. When a zener diode is operating in the breakdown region, an increase in zener voltage causes the (a) diode's resistance to decrease, (b) reverse current to decrease, (c) forward current to increase.
7. The PNPN unit of Figure 6 is forward biased when (a) the inner P-region is positive with respect to the inner N-region, (b) the inner P-region is positive with respect to the outer P-region, (c) the outer N-region is negative with respect to the outer P-region.
8. When a PNPN device is in the forward blocking region, forward current is held to a small value by the reverse biased middle junction. True or False?
9. A PNPN device remains in its forward high conduction region until the forward current falls below a value called the (a) breakdown current, (b) holding current.

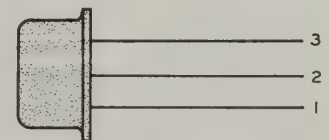


A trigistor is shown at the left in this photograph. On the small board at the right, one trigistor, one capacitor, three resistors and two diodes form a circuit capable of serving the same function as the two transistors and other components on the large board.

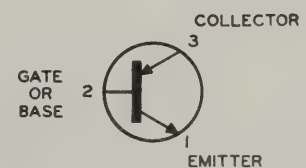
Courtesy Solid State Products, Inc.

Various applications of PNP devices require a wide range of current-carrying abilities. Units designed for low current often have the general appearance of Figure 10A. Figure 10B gives the circuit symbol employed for these units.

The outer N-region is called the emitter, the outer P-region the collector, and the internal P-region the gate. Sometimes this latter region is referred to as the base. As indicated, lead No. 1 is the emitter lead, No. 2 is the gate lead, and No. 3 the collector lead. Various trade names of these units have included TRIGISTOR, TRANSWITCH, and DYNAQUAD.



A



B

Figure 10

Higher current PNP devices are normally called **CONTROLLED RECTIFIERS**. Most often, the semiconductor used is silicon, in which case the unit is usually referred to as a **SILICON CONTROLLED RECTIFIER (SCR)**. The appearance of a typical unit is shown in Figure 11A. In the SCR, the outer N-region is called the cathode, the inner P-region the gate, and the outer P-region the anode. Figure 11B gives the circuit symbol used for the SCR, while the internal construction is shown by the drawing of Figure 11C.

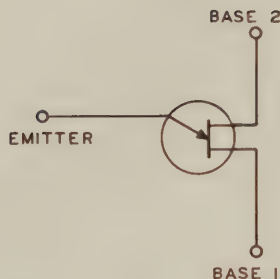


Figure 12

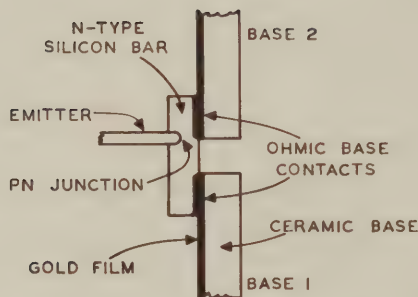


Figure 13

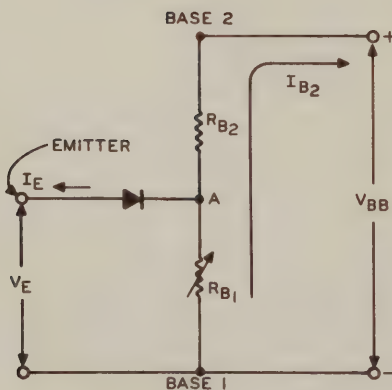


Figure 14

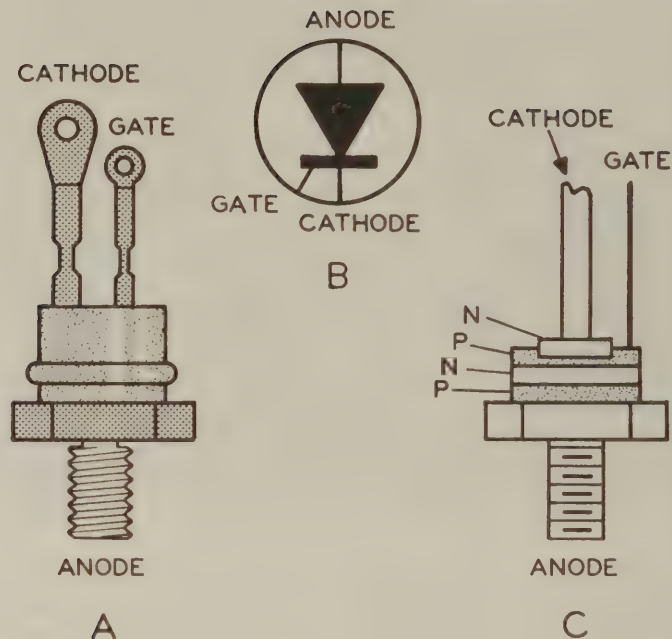


Figure 11

UNIUNCTION TRANSISTORS

The **UNIUNCTION TRANSISTOR (UJT)** is a unique semiconductor control device. Although it is called a transistor, its characteristics are vastly different from those of a conventional junction transistor. The conventional transistor normally operates as a linear amplifier, whereas a UJT has the ability to "fire" or switch from low to high conduction at certain voltage levels, and is thus used in control applications.

Figure 12 shows the symbol for a unijunction transistor. The electrodes are labeled EMITTER, BASE 1 and BASE 2. Figure 13 shows the physical construction of one type of unijunction transistor. The N-type silicon bar is mounted on a ceramic base. Ohmic contacts are made to each end of the bar. These form the base leads. The emitter PN junction is formed where the emitter lead connects to the N-type silicon bar.

Figure 14 shows the simplified equivalent circuit of the unijunction transistor. R_{B1} and R_{B2} are the equivalent resistances of base 1 and base 2. The diode in the emitter leg represents the PN junction shown in Figure 13.

Normally, base 2 is made more positive than base 1. The interbase voltage, V_{BB} , establishes a current I_{B2} as shown in Figure 14. This current through R_{B1} and R_{B2} produces an internal voltage at point A (across R_{B1}) which makes the cathode end of the "emitter diode" positive. Emitter voltage V_E

must exceed this internal voltage before appreciable current will flow in the emitter circuit.

Figure 15 shows the emitter characteristics for a unijunction transistor. Note that emitter voltage (V_E) is plotted vertically, and emitter current (I_E) is plotted horizontally. As emitter voltage is increased toward the value of the internal voltage (the voltage at A in Figure 14), determined by internal resistances, emitter current remains at a low cutoff value. Once the emitter voltage reaches the value of the internal voltage, the emitter PN junction becomes forward biased and emitter current increases to a value above the cutoff value. The point where conduction occurs is called the **PEAK POINT**. The value of emitter voltage necessary to produce conduction is called the **PEAK POINT VOLTAGE (V_P)**.

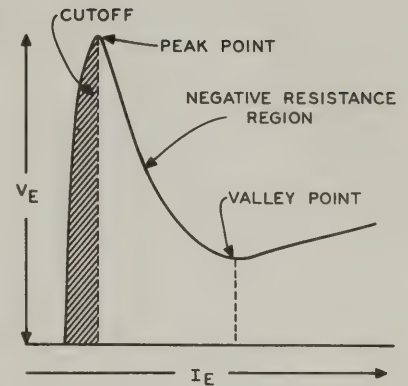


Figure 15

A very unusual action now occurs. The emitter-to-base 1 resistance, R_{B1} in Figure 14, begins to decrease faster than the emitter current increases. This gives rise to a **NEGATIVE RESISTANCE** characteristic as shown by a decrease in emitter voltage as emitter current increases. This negative resistance continues until the **VALLEY POINT** of minimum emitter voltage is reached. Beyond this point, emitter voltage increases in a normal manner as emitter current increases.

In summary, the base 1 resistance is very high at low emitter voltages. It decreases very suddenly as the emitter voltage reaches a level sufficient to cause forward emitter bias. At this point the resistance appears to become negative, because emitter voltage decreases as emitter current increases.

The negative resistance characteristics of the unijunction transistor makes it especially suited for switching applications. It can itself be used as a switch or, as is often the case, it can be used to control other semiconductor switching devices. A unijunction, operating as a relaxation oscillator, is often used for such control purposes. Figure 16 shows a unijunction relaxation oscillator circuit. Figure 17 shows the important circuit waveshapes.

Let's examine the operation of the circuit. The emitter PN junction must be forward biased for the unijunction emitter to conduct. In this circuit, R_3 is small and the voltage across C_1 is approximately the emitter voltage. As C_1 charges, the emitter gradually becomes more positive. This exponential rise can be seen as the rising portion of each cycle in Figure 17A.

When the emitter voltage reaches the peak voltage, V_P , the emitter PN junction becomes forward biased and heavy conduction takes place in the emitter circuit. As emitter circuit current increases, the emitter voltage de-

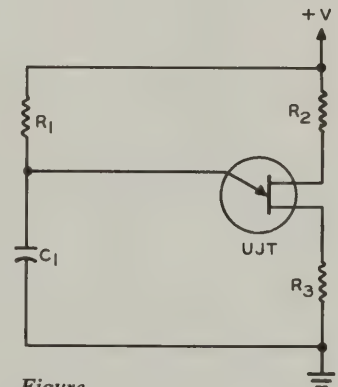


Figure 16

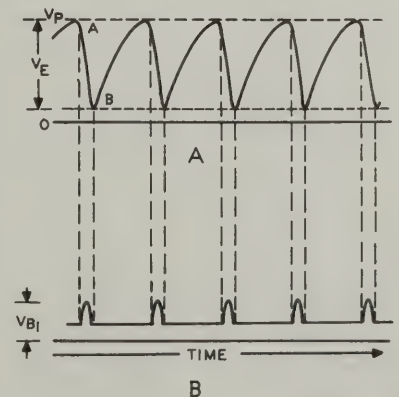
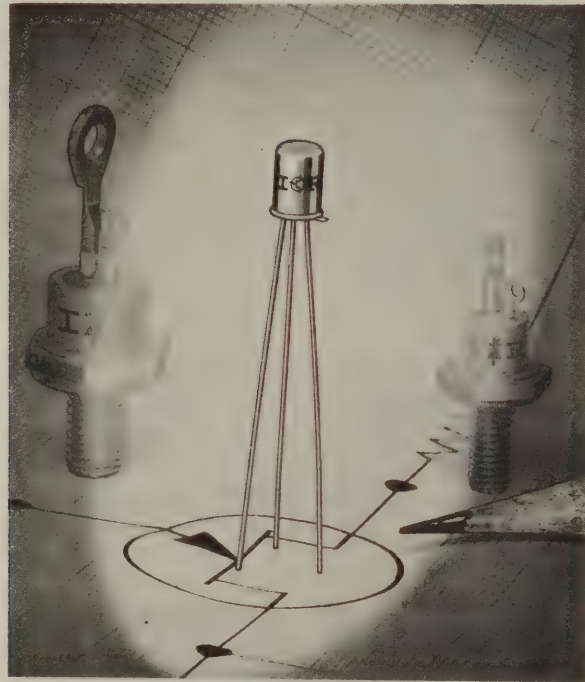


Figure 17

creases due to the negative resistance characteristic mentioned earlier. The net action is the very rapid discharge of C_1 . This is shown by the rapid drop in the emitter voltage waveform between points A and B in Figure 17A.



A unijunction transistor designed for use as a switch to control the current in the gate circuit of a controlled rectifier.

Courtesy International Rectifier Corp.

The emitter voltage decreases rapidly until the emitter PN junction becomes reverse biased. At this point, emitter current stops and the emitter resistance reverts to its high (reverse biased) value. With I_E cut off, the capacitor begins to charge and the cycle repeats.

Each time emitter conduction takes place, a heavy pulse of current passes through R_3 . As a result, a pulse of voltage is developed across R_3 as shown in Figure 17B. This pulse can be used to trigger another control device.

The frequency of the pulses developed by the circuit of Figure 16 depends primarily upon the values of R_1 and C_1 , for any given UJT. These two components determine the charging rate of the capacitor. This, in turn, determines the rate at which the emitter voltage rises to the conduction value.

The actual voltage for conduction is determined by the **INTRINSIC STAND-OFF RATIO** of the UJT. For example, if this ratio is $\frac{1}{2}$, the emitter conducts when it reaches $\frac{1}{2}$ the base-to-base voltage. If the intrinsic stand-off ratio of the UJT of Figure 16 is $\frac{1}{2}$, and the base-to-base voltage is 10 volts, the emitter will conduct when it reaches a voltage 5 volts more positive than base 1 ($\frac{1}{2} \times 10 = 5$). This circuit is reasonably stable against changes in the supply voltage V . If the supply voltage increases, the base-to-base voltage increases also. The stand-off ratio remains comparatively constant so the conduction voltage needed at the emitter also increases. However, C_1 charges a little more rapidly through R_1 because of the rise of voltage V . Therefore, the voltage across C_1 reaches V_P in the same amount of time as it did before. As a result, the frequency remains about the same. This is an advantage over some types of RC oscillators, in which the operating frequency is highly dependent on supply voltage.

ELECTRONIC SWITCHING

All electrically and electronically operated equipment must be switched on to begin operation and off to stop. In addition, many devices require frequent or continual switching in the internal circuits as part of their normal operation. Various mechanical and electromechanical switches are used. However, frictional wear and inertia limit the usefulness of these devices for either high speed or continuous operation. Such applications normally employ an electronic switch of some type.

Any electronic device that can change back and forth between states of low and high conduction can serve as a switch. The general requirements of a good switch are:

- (1) The state of little conduction, or the "off" state, should provide practically an open circuit condition.
- (2) The "on" state should provide practically a short circuit through the switch device.
- (3) The amount of power needed to hold the switch should be small.
- (4) The switch should open or close very rapidly.

The last requirement is met most easily by electronic switches since they do not have moving parts. The PNP devices meet all of these requirements very well and therefore are commonly employed as electronic switches.

BASIC PNPN SWITCHING CIRCUITS

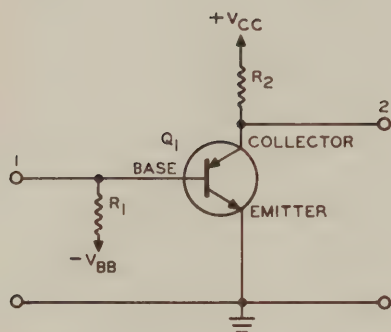
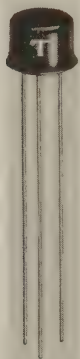


Figure 18



The transwitch is a low power PNPN device for switching applications.
Courtesy Transistron Electronic Corp.

Figure 18 shows a low-power PNPN device employed as a switch. The purpose of this circuit is to switch the output terminal 2 alternately between a high and a low positive voltage. The $+V_{CC}$ voltage applies forward bias to the collector of Q_1 . However, this voltage is less than the breakover voltage. Therefore, the reverse bias on the middle junction holds Q_1 in its off state.

The nonconductive state is maintained by voltage $-V_{BB}$ also. This voltage is applied through R_1 to the base, and reverse biases the base-emitter junction. With Q_1 nonconductive there is no current in resistor R_2 and no voltage across R_2 . Hence, terminal 2 and the Q_1 collector are at the positive potential of $+V_{CC}$.

Q_1 is switched to its on state by applying a positive voltage pulse to terminal 1. This pulse forward biases the base-emitter junction, producing base current. As the base current rises, it causes the breakover voltage of Q_1 to decrease. When V_{B0} reduces to the value of $+V_{CC}$, forward breakdown occurs and Q_1 shifts into its conductive or on state.

With Q_1 conducting there is a voltage across R_2 equal to the product of the collector current and the resistance of R_2 . Also, when Q_1 is conducting, its resistance is very much lower than that of R_2 . Therefore, the voltage across Q_1 is practically zero, while that across R_2 is practically equal to the entire $+V_{CC}$ supply voltage. Thus, when Q_1 turns on, the potential of terminal 2 falls to nearly zero.

Figure 19 illustrates a basic method of employing a silicon controlled rectifier as a switch. Here, SCR_1 is in series with the load. SCR_1 acts like an open circuit when it is off, and there is no current in the circuit made up of B_1 , SCR_1 and the load.

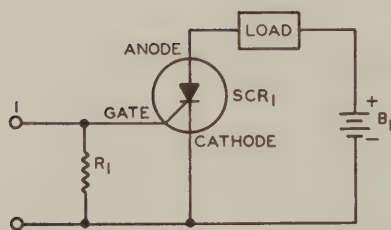


Figure 19

The potential of the positive terminal of B_1 is applied through the load to the anode of SCR_1 . The anode voltage applied by the battery forward biases SCR_1 . However, with no gate current the breakover voltage is greater than the anode voltage, so SCR_1 remains off.

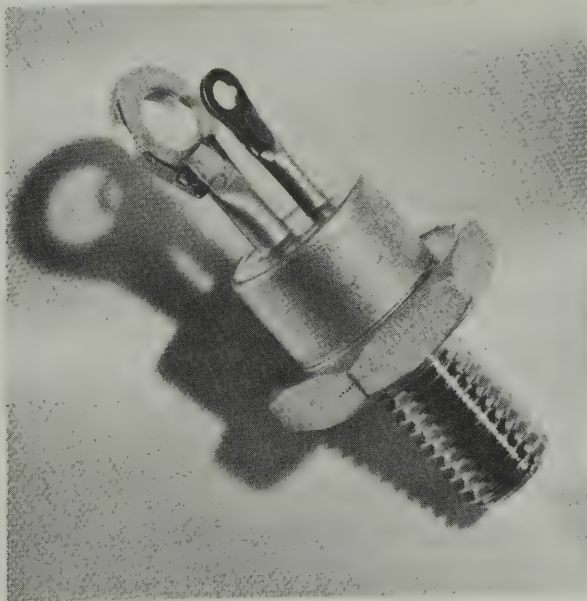
To turn SCR_1 on, a voltage applied to terminal 1 must produce gate current large enough to reduce the breakover voltage to the value of the B_1 voltage. SCR_1 then becomes conductive between cathode and anode, permitting B_1 to supply current to the load. SCR_1 remains on until the anode current is reduced below the holding current level.

TURN-ON AND TURN-OFF METHODS

PNPN devices can be turned on by applying a voltage that produces gate (or base) current. Such a current is produced in Figure 19 when the applied turn-on or trigger voltage makes the gate positive with respect to the cathode. The gate current must be increased to a value that makes the forward break-over voltage equal to the anode-cathode voltage. It is also called the "gate current to trigger", I_{GT} , or "gate current to fire", I_{GF} .

Some four-layer devices can be turned off by applying a negative pulse to the gate. These devices are often referred to as GATE CONTROLLED SWITCHES, GATE TURN-OFF CONTROLLED RECTIFIERS, etc. Silicon controlled rectifiers (SCR's) cannot be turned off by applying a pulse to the gate. The anode voltage must be reduced for the gate to regain control. This action is similar to that in a gas filled tube called a thyatron. For this reason an SCR is often called a "solid state thyatron."

Either dc or ac voltages can be employed to operate the PNPN devices. Figure 20 shows an SCR circuit in which both the gate trigger signal and



The terminals are clearly illustrated in this photograph of a silicon controlled rectifier. The mounting stud is the anode terminal, the large ring is the cathode terminal, and the small ring is the gate terminal.

*Courtesy General Electric Co.
Semiconductor Products Dept.*

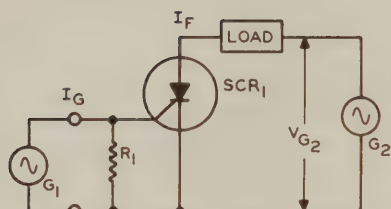
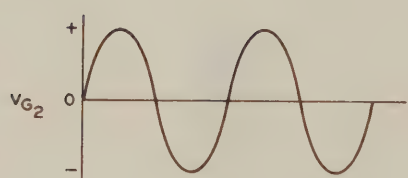


Figure 20

the anode supply voltage are ac voltages. Sources G_1 and G_2 supply voltages that are of the same frequency. When these voltages are in phase, the gate current I_G produced by G_1 is in phase with the G_2 voltage, V_{G2} , as shown in Figure 21.

The positive rise of V_{G2} makes the anode positive with respect to the cathode. However, forward anode current I_F remains at zero until the gate reaches the level I_{GT} . At this point, forward breakdown occurs, and I_F rises suddenly to a value determined by V_{G2} and the resistance of the load. After this point, the gate no longer has control, and I_F varies with V_{G2} .



After the positive peak, V_{G2} falls toward zero, causing I_F to fall also. When I_F drops below the holding current level, I_H , the SCR quickly switches back to its off condition, making I_F drop immediately to zero. I_F remains at zero until I_G again reaches the gate current to trigger level, I_{GT} , in the next cycle.

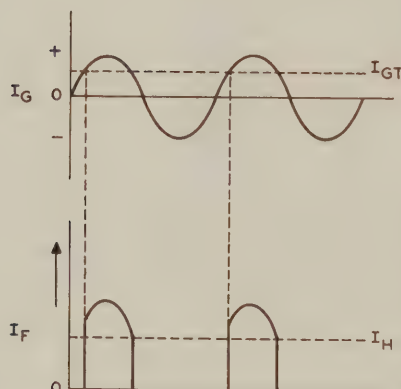


Figure 21

This operation supplies load current in the form of pulses. The length of each pulse depends upon the interval of time between the instant at which I_G reaches the I_{GT} level and the instant at which I_F falls to the I_H level. The pulse length can be changed by shifting the phase of I_G with respect to V_{G2} . The SCR can be made to turn on or "fire" at any point in the positive alternation of V_{G2} at which this voltage is large enough to make I_F greater than I_H . Another requirement for proper operation of this circuit is that the peak value of V_{G2} must be less than either the forward or reverse breakdown voltage of the SCR.

As described earlier, a normal SCR cannot be turned off by applying a pulse to the gate. Instead, the anode voltage must be reduced to reduce the forward current for the gate to regain control.

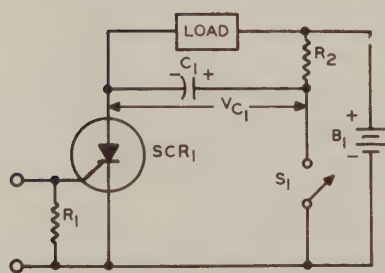


Figure 22

With an ac anode voltage as shown in Figure 20, control is regained each cycle since the anode voltage reverses polarity. However, other methods of obtaining turn-off can also be used.

Figure 22 shows a method of turn-off referred to as PARALLEL TURN-OFF. When SCR_1 is conducting, its anode current passes through the load. Capacitor C_1 charges through R_2 to the voltage across the load. The polarity of the C_1 charge is indicated in the figure. When it is desired to make SCR_1 turn off, switch S_1 is closed. This connects C_1 directly across SCR_1 so that V_{C1} makes the anode negative with respect to the cathode. This reverse bias turns SCR_1 off, and C_1 discharges through the load and B_1 .

The following Practice Exercise questions cover the subjects which you have just studied. They are:

GATE CONTROLS

UNIJUNCTION TRANSISTORS

ELECTRONIC SWITCHING

BASIC PNP SWITCHING CIRCUITS

10. The larger the gate current in a PNP unit, the higher the breakover voltage, V_{BO} . True or False?
11. In the silicon controlled rectifier, what name is given to the inner P-region?
12. Draw the symbol for a unijunction transistor and label the electrodes.
13. A unijunction transistor has two PN junctions. True or False?
14. For normal operation of the unijunction transistor shown in Figure 12, base 2 is made positive with respect to base 1. True or False?
15. Emitter conduction occurs in the unijunction transistor when the emitter junction is (a) forward biased, (b) reverse biased.
16. In the circuit of Figure 16, C_1 charges to the point where the UJT (a) conducts, (b) cuts off.
17. What determines the frequency of the pulses developed by the circuit of Figure 16?
18. In the circuit of Figure 18, PNP unit Q_1 can be switched on by applying (a) a negative pulse to the collector, (b) a negative pulse to the base, (c) a positive pulse to the base.
19. In Figure 19, SCR_1 switches into the highly conductive region when the gate current is large enough to reduce the breakover voltage to a value equal to the B_1 voltage. True or False?

Figure 23 shows an arrangement known as SERIES TURN-OFF. When SCR_1 is nonconductive, any charge on C_1 leaks off through the load, which is in the cathode circuit. As soon as SCR_1 is triggered on, C_1 begins to charge through SCR_1 and L_1 . As C_1 charges, point A becomes more and more positive. Since point A connects to the cathode of SCR_1 , the rising voltage across C_1 carries the cathode potential toward that of the anode. This decreases the forward bias on SCR_1 , and the anode current begins to decrease. The field about L_1 collapses, inducing a voltage that tends to maintain the current. This induced voltage aids that of B_1 , causing C_1 to charge to approximately twice the B_1 voltage.

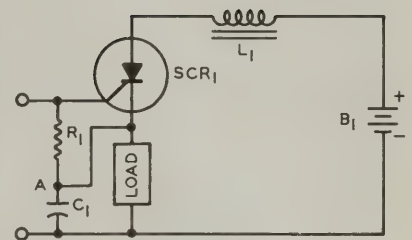


Figure 23

Finally, the anode current drops below the holding current level, and SCR_1 turns off. With no anode current, there is no voltage across L_1 , and the anode has the potential of the positive terminal of B_1 . At this point, the large C_1 voltage makes the cathode of SCR_1 more positive than the anode. This reverse bias prevents SCR_1 from being triggered on until C_1 discharges through the load to a low value.

SCR SWITCHING CIRCUITS

Controlled rectifiers are widely used in switching applications. Figure 24 shows a pair of controlled rectifiers used in an ac switching circuit. They provide a means of opening or closing the path through which the ac source supplies a large alternating current to the load. SCR_1 and SCR_2 have the anode of one connected to the cathode of the other. Together they act as a switch between the load and ac input terminal 2. Hence, current can flow in the load only when the SCR switch circuit is closed.

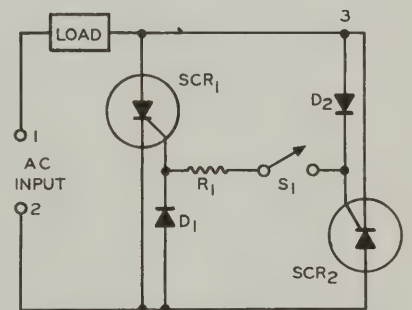


Figure 24

As shown, the rectifiers and diodes are connected so that when the ac input voltage makes the anodes of SCR_1 and D_2 positive, those of SCR_2 and D_1 are negative. With switch S_1 open, a positive voltage is applied to the gate of a rectifier only when this rectifier is reverse biased. Hence, no anode current is produced.

For example, when the ac input voltage makes terminal 1 positive, this positive potential is applied through the load to the anodes of SCR_1 and D_2 . At this time, terminal 2 applies a negative potential to the anodes of SCR_2 and D_1 . Therefore, D_1 is nonconductive so that there is no voltage applied to the gate of SCR_1 . Although forward biased, SCR_1 remains in its off state. With its anode positive, D_2 is conductive, and therefore connects point 3 to the gate of SCR_2 . However, as the anode of SCR_2 is negative with respect to the cathode, this rectifier remains in its off state also.

A similar set of conditions exists when the ac input voltage makes terminal 1 negative and terminal 2 positive. Then, SCR_1 is off because it is reverse biased, and SCR_2 is off because its gate is disconnected by nonconductive diode D_2 .

The controlled rectifier switch circuit is made conductive for load current by closing switch S_1 . Then, when the ac input makes terminal 1 positive, SCR_1 is forward biased as before, and a positive potential is applied through D_2 , S_1 and R_1 to the gate of SCR_1 . This produces gate current in the form of electron flow from input terminal 2 to the cathode of SCR_1 , through the cathode-gate junction to the gate, through R_1 , S_1 and D_2 to point 3, and through the load to input terminal 1. When the gate current reaches the I_{GT} level, SCR_1 turns on and conducts a pulse of anode current which flows through the load. SCR_2 is nonconductive because it is reverse biased at this time.

Likewise, with S_1 closed, when the ac input makes terminal 2 positive, a positive potential is applied through D_1 , R_1 and S_1 to the gate of SCR_2 . Gate current consists of electron flow from input terminal 1 to point 3, through the cathode-gate junction of SCR_2 , and through S_1 , R_1 and D_1 to input terminal 2. As a result, SCR_2 turns on and conducts a pulse of anode current which passes through the load.

SCR_1 is reverse biased and therefore turned off at this time. The conduction of SCR_2 produces electron flow through the load in the opposite direction

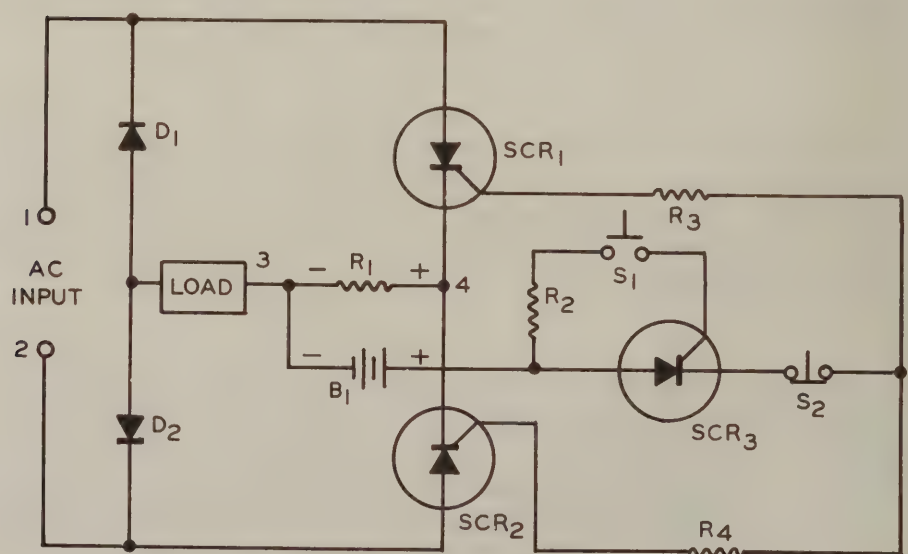


Figure 25

to the flow produced by the conduction of SCR_2 . Therefore, the load receives alternating current. This current has the form of pulses whose length depends upon the time that the respective controlled rectifiers remain in their on states.

Figure 25 shows a circuit in which two diodes and two controlled rectifiers form a bridge rectifier which supplies direct current to a load from an ac source. When input terminal 1 is positive, the load current path consists of D_2 , the load, R_1 and SCR_1 . When input terminal 2 is positive, D_1 , the load, R_1 and SCR_2 form the load current path.

Controlled rectifiers SCR_1 and SCR_2 are connected so that the ac input voltage makes the anode of one SCR positive when it makes the anode of the other SCR negative. As the applied input voltage rises from zero, the rectifier with the positive anode turns on at some point, provided gate current is supplied to it from the SCR_3 circuit.

Low-voltage source B_1 applies a positive voltage to the anode of SCR_3 , and also to the gate when push-button switch S_1 is closed. SCR_3 remains on after being triggered on by a momentary positive voltage applied to its gate. With SCR_3 turned on, its anode current passes through B_1 , R_1 , the cathode-gate junctions of SCR_1 and SCR_2 , resistors R_3 and R_4 , and normally-closed switch S_2 . Thus, the anode current of SCR_3 serves as gate current for SCR_1 and SCR_2 .

SCR_3 is employed as a protective device. The protective action is as follows: load current flows through resistor R_1 from the load to the cathodes of SCR_1 and SCR_2 . The resulting dc voltage across R_1 makes point 4 positive with respect to point 3, as indicated. This voltage varies directly with the load current. The cathode of SCR_3 has very nearly the same potential as point 4 because the voltages across R_3 and R_4 and the cathode-gate junctions of SCR_1 and SCR_2 are very small. If an increase in load current makes the voltage across R_1 sufficiently greater than the B_1 voltage, the cathode of SCR_3 becomes more positive than the anode. This reverse bias turns SCR_3 off. With no gate current supplied to them, SCR_1 and SCR_2 cannot conduct anode current. Thus, all three SCR's become nonconductive, and the load current is interrupted. When desired, the SCR circuit can be turned off also by momentarily opening push-button switch S_2 to interrupt the SCR_3 anode current.

SCR CONTROL CIRCUITS

SCR's can be used to vary the amount of power applied to a load by controlling the length of time the SCR conducts during each cycle. Typical

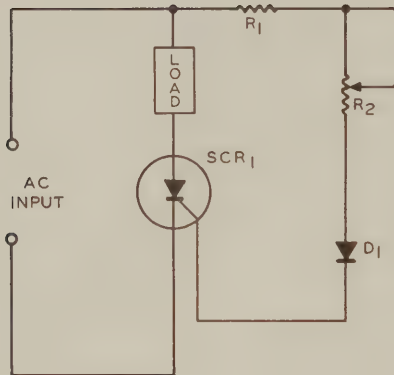


Figure 26

applications would include motor speed controls, lighting level controls, variable power supplies, etc. Figure 26 shows an SCR circuit using a simple amplitude control. The waveform of the voltage appearing across the load for various settings of R_2 is shown in Figure 27. The shaded part of the waveform represents the time during which the SCR conducts. Notice that control is limited to about half of the positive alternation. That is, conduction can be varied from about 0° to 90° .

The circuit of Figure 26 is basically a controlled half wave rectifier circuit. SCR_1 and the load form a series circuit connected to the ac supply. Like a regular diode, SCR_1 can conduct only when its anode is positive with respect to its cathode. Thus, power is supplied to the load on only one alternation. By employing two SCR's, as shown earlier, power can be supplied to the load on both alternations to obtain full wave operation.

The triggering circuit of Figure 26 consists of R_1 , R_2 , and D_1 . This circuit applies the positive pulse of the ac voltage appearing at the SCR_1 anode to the gate. Diode D_1 prevents the gate from being driven negative. R_2 controls the value of gate current, and R_1 limits gate current to a safe value when R_2 is set to minimum resistance.

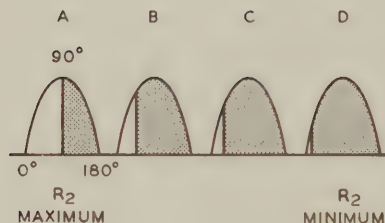


Figure 27

Suppose R_2 is set to its maximum value and the ac supply is on its positive alternation, making the anode of SCR_1 positive with respect to the cathode. As the ac supply increases toward its maximum value, the resistance of R_2 is high enough to hold the gate current below the value required to cause conduction. If the total resistance of R_1 and R_2 were a high enough value, SCR_1 would never conduct. However, suppose the values of R_1 and R_2 are such that the required value of gate current is reached at the peak of the positive ac supply alternation. SCR_1 conducts at this point and continues to conduct for the remainder of the alternation as shown by the shaded area on Figure 27A.

To turn off an SCR once it is conducting, the anode current must be reduced below the holding current level or the anode voltage must be reversed as described earlier. This is automatically obtained with an ac supply because the alternating voltage drops to zero twice each cycle. As the alternating voltage drops toward zero, the current through the load decreases. Sometime before the end of the alternation, the load current will drop to the holding level and the SCR will turn off. On the next alternation, the polarity of the ac anode voltage is reversed, keeping SCR_1 cut off. Thus, the gate can regain control once each cycle.

Now suppose the value of R_2 is decreased. This allows the gate current to reach the required value earlier in the alternation, and thus SCR_1 conducts earlier as shown in Figure 27B. Further reducing the value of R_2 causes

SCR₁ to conduct still earlier as shown in Figure 27C. Finally, SCR₁ can be made to conduct for almost the full half-cycle as shown in Figure 27D.

Increasing the conduction time supplies power to the load for a greater portion of each cycle, thus increasing total load power. One disadvantage of this circuit is that control can be had for only about one-half of each alternation, or one-quarter of each cycle (0 to 90°).

A greater range of control can be obtained by using a phase shift circuit to trigger the gate. Figure 28 shows a half wave SCR circuit employing an RC network in the trigger circuit. The RC network permits variation of SCR conduction over almost a complete half-cycle (180°).

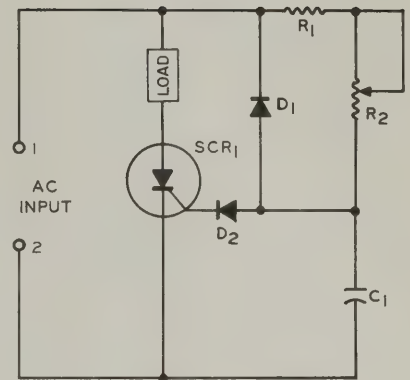


Figure 28

The triggering circuit consists of diodes D₁ and D₂, resistors R₁ and R₂, and capacitor C₁. SCR₁ is connected in series with the load as in Figure 26. Note however, that SCR₁ and the load are across the source. This circuit, as in the circuit of Figure 26, applies the ac voltages to the anode of the SCR and to the gate trigger circuit. D₂ permits only positive pulses to reach the gate of SCR₁. D₁ discharges C₁ each cycle (actually charges it to the peak of the negative alternation) so that the triggering is uniform from cycle to cycle.



By controlling the length of time the SCR conducts during each cycle, the amount of power supplied to the load can be varied. This effect is used in the motor control unit shown which provides a wide range of speed and power control.

Courtesy Seco Electronics, Inc.

The gate voltage in Figure 28 is developed across C_1 . The combination of R_1 , R_2 and C_1 makes up a phase shift circuit which provides a lagging voltage compared to that at the SCR anode. Thus, the voltage increase across C_1 occurs later in the cycle than that at the SCR anode.

In Figure 28, the maximum value of R_2 produces the maximum amount of lag or phase shift and results in firing late in the alternation as shown by the shaded area in Figure 29A. Decreasing the value of R_2 reduces the phase shift and causes conduction of SCR_1 earlier in the cycle as shown in Figure 29B. Further reduction in the value of R_2 causes SCR_1 to conduct still earlier as shown in Figure 29C and 29D. Depending on the values of R_2 and C_1 , you can vary the load current over most of the positive half cycle as shown in Figure 29.

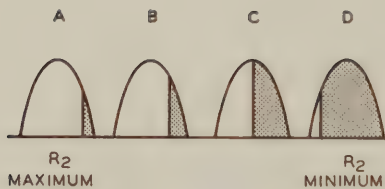


Figure 29

Pulse triggering of an SCR can give a wide range of positive control. One of the uses for the unijunction transistor is to trigger an SCR by means of pulses applied to the gate. A circuit including a UJT is shown in Figure 30. D_1 rectifies the ac input voltage to provide a pulsating dc operating voltage for the unijunction trigger circuit. In this circuit the gate voltage is developed across R_3 . In order for R_3 to develop enough voltage to trigger SCR_1 , the unijunction must conduct. Capacitor C_1 along with resistors R_1 and R_2 make up an RC circuit which controls the conduction of UJT_1 . When C_1 charges to a high enough voltage, the emitter circuit of unijunction UJT_1 conducts heavily. This current produces a voltage pulse across R_3 which triggers SCR_1 .

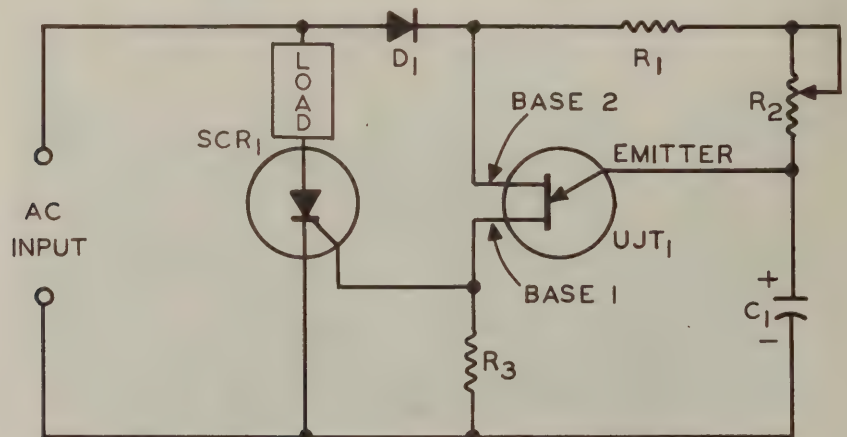


Figure 30

Let's take a closer look at the operation of Figure 30. On the positive alternation of the input cycle, the anode of SCR_1 is made positive with respect

to the cathode. Likewise, base 2 of UJT_1 is made positive with respect to base 1. At the same time, C_1 begins to charge at a rate determined by the value of C_1 , R_1 and R_2 . When the voltage across C_1 increases to a large enough value, the emitter of UJT_1 conducts. Emitter current from C_1 develops a voltage pulse across R_3 . This voltage pulse is applied to the gate of SCR_1 . Assuming that the pulse is of sufficient amplitude, SCR_1 conducts.

With R_2 set to a high value, C_1 charges slowly to cause the UJT emitter conduction late in the cycle. In turn, SCR_1 conducts later in the cycle. Reducing the value of R_2 causes UJT_1 to conduct earlier, resulting in an increase in load power.

CONTROLLED RECTIFIER RATINGS

No single semiconductor device is perfectly suited to every possible use. For example, although useful in many applications, transistors are limited in their current and voltage ratings. Consequently, they cannot be employed in the many applications that require the control of very large currents at several hundred volts. However, this need can be filled by the controlled rectifier. The smaller PNP devices of the trigistor type have the advantage of compactness. Where they can be used, often one of them performs the functions that would require two transistors.

An important advantage of the controlled rectifier in switching applications is in the amount of trigger signal current needed to produce full conduction. In many power transistors, for example, it is necessary to supply up to 500 milliamperes to the base to make the transistor conduct 5 amperes from emitter to collector. A typical controlled rectifier that can conduct an anode current up to 5 amperes can be triggered on by a gate current of only 10 mA or less.

Other important features of PNP devices are high efficiency (about 99% compared with about 60% for transistors), very low voltage across the device when it is conducting, very fast switching (on and off), small size and weight, and extreme ruggedness when operated within their normal limits.

Various ratings of PNP devices include the peak reverse voltage, rms forward current, peak gate current and peak gate voltage (forward and reverse), and minimum forward breakover voltage.

The peak reverse voltage (PRV), is the maximum allowable reverse voltage that may be applied repeatedly at short intervals across the PNP device. Typical PRV ratings range from 25 to over 1000 volts.

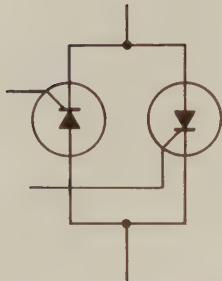


Figure 31

The rms forward current is the maximum allowable rms value of the forward anode current pulses, or the maximum allowable direct anode current. Ratings range from 200 mA to well over 200 amperes for various types.

The peak gate current and voltage are the maximum allowable gate current and voltage, respectively. Typical ratings are from 0.1 to 4 amperes.

Finally, the minimum forward breakover voltage is the lowest positive anode voltage at which a controlled rectifier will switch into the conductive state with the gate circuit open.

TRIAC CONTROL CIRCUITS

A special semiconductor device called a **TRIAC** is similar to two SCR's connected as shown in Figure 31. This device can conduct in both directions, when the proper gating signal is applied, which permits it to control alternating current. SCR's can also be used, but this requires two units as shown earlier in Figure 24.

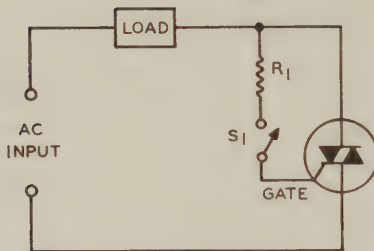


Figure 32

A basic triac circuit is shown in Figure 32. Notice that the symbol used for the triac is different from that of an SCR. The single control element is still called the gate. The other leads cannot be labeled anode and cathode, however, since the device can conduct in both directions. The triac is connected in series with the load in a manner just like the SCRs in the previous circuits.

When S_1 is closed, gate current is supplied to the triac on each alternation through resistor R_1 which limits the gate current to the proper value. Thus, the triac conducts on both alternations of the input. Since it takes a minute amount of voltage to start the triac conducting, there is a small portion of each alternation during which there is no conduction. However, for practical purposes we assume the triac conducts for the full input cycle.

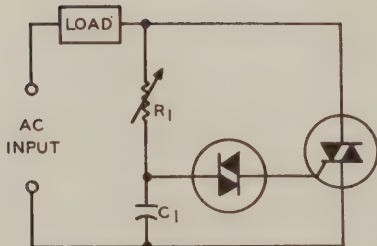


Figure 33

The primary use of the circuit of Figure 32 is on-off control of heavy current loads. Switch S_1 breaks and makes only the light current gate circuit, whereas the triac itself acts as the switching element for the heavy current load circuit. Hence, a small switch S_1 can control large currents by means of the triac.

Figure 33 shows how phase control can be applied to a triac circuit. The phase shift network consists of variable resistor R_1 and capacitor C_1 . In the gate circuit of the triac is another special semiconductor device called a **DIAC**. A diac might be thought of as two zener diodes connected back-to-

back. The diac conducts, in either direction, whenever a voltage greater than its breakover voltage is applied.

When the voltage across C_1 reaches the diac breakover voltage, the diac conducts, allowing C_1 to discharge into the triac gate. This pulse triggers the triac into conduction for the remainder of that half cycle. The diac can conduct on both alternations; hence it applies a pulse to the triac during each alternation. In turn, the triac conducts during both alternations and thus permits alternating current to flow in the load.

Varying the value of R_1 controls the phase shift or lag of voltage across C_1 . With R_1 set to a low value, the diac conducts early in each alternation, thus allowing the triac to conduct early. The load power is relatively high under these conditions. Increasing the value of R_1 delays conduction of the triac in each alternation. Thus, current is permitted in the load over a smaller part of each alternation. This results in reduction of ac power to the load.

SUMMARY

Practically every branch of the electronics field employs various devices for controlling voltage and current, changing ac to dc, and changing voltages to values other than those provided by the power lines. Semiconductor devices of this type include two, three, and four-layer units with power ratings ranging from very low to very high values.

Two-layer devices are PN junction diodes and rectifiers. Zener diodes and avalanche diodes provide control of voltages. Conversion from ac to dc is accomplished with semiconductor rectifiers.

The four-layer devices comprise two categories. All block forward current until the gate current or the forward bias across the entire unit brings the unit to its breakover point. One category is made up of the low-power PNP devices, such as trigistors, transwitches, and dynaquads. The other category is composed of the high-power PNP devices which are more commonly called controlled rectifiers. Once an increase in gate current has caused one of these units to conduct, the unit can no longer be controlled by the gate until the conduction is stopped by reducing the anode current below the holding current level.

Proportional control of power applied to a load can be obtained by controlling the length of time the controlled rectifier conducts during each supply cycle. Simple resistive control circuits can provide only 90° of control each alternation. Phase control circuitry, on the other hand, can provide control for almost a full alternation, 180°.

Unijunction transistors are widely used to control the conduction of controlled rectifiers. The unijunction transistor is different from regular transistors in that it is basically a switching device.

The triac is similar to two controlled rectifiers connected in inverse parallel. This device can conduct in both directions and can thus control ac power.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

TURN-ON AND TURN-OFF METHODS

SCR SWITCHING CIRCUITS

SCR CONTROL CIRCUITS

CONTROLLED RECTIFIER RATINGS

TRIAC CONTROL CIRCUITS

20. In the circuit of Figure 20, shifting the phase of I_G with respect to V_{G_2} has what effect on the pulses of anode current I_F ?
21. In the circuit of Figure 25, SCR_1 and SCR_2 gate currents can flow only when (a) SCR_3 is on, (b) SCR_3 is off, (c) S_2 is open, (d) S_1 is open.
22. How can the amount of power applied to a load be varied with an SCR?
23. With R_2 set to maximum resistance in Figure 26, the SCR conducts for (a) almost all of the positive alternation, (b) only a portion of the positive alternation.
24. Maximum power is supplied to the load when R_2 is set to (a) maximum resistance, (b) minimum resistance.
25. Which component in Figure 28 develops the SCR gate voltage?
26. SCR conduction in Figure 28 can be varied over almost (a) half of the positive alternation, (b) the full positive alternation.
27. If C_1 in Figure 28 takes a long time to charge, the SCR conducts (a) earlier during the positive supply alternation, (b) later during the positive supply alternations.
28. Maximum power is supplied to the load in Figure 28 when R_2 is set to maximum. True or False?
29. The gate voltage for the circuit of Figure 30 is developed across (a) R_1 , (b) R_3 , (c) C_1 .
30. A gate voltage is developed across R_3 in Figure 30 whenever UJT_1 conducts. True or False?
31. SCR_1 in Figure 30 begins to conduct when UJT_1 (a) conducts, (b) is cut off.

32. In Figure 30, what controls the point at which UJT_1 conducts during the positive supply alternation?
33. The device best suited to the control of very large currents at high voltages is the (a) power transistor, (b) controlled rectifier, (c) diode vacuum tube.
34. Compared with the controlled rectifier, full conduction in a power transistor requires a larger trigger signal current. True or False?
35. Compared with the transistor, the controlled rectifier is (a) more efficient, (b) less efficient.
36. In Figure 32, the triac can conduct on (a) one alternation of the ac supply, (b) both alternations of the ac supply.
37. The diac in Figure 33 can be thought of as two zeners connected back-to-back. True or False?
38. Increasing the value of R_1 in Figure 33 causes a reduction in the ac power to the load. True or False?

IMPORTANT DEFINITIONS

AVALANCHE DIODE — See BREAKDOWN DIODE.

BREAKDOWN DIODE — A diode designed to operate in the reverse breakdown region. Also called an AVALANCHE DIODE or ZENER DIODE.

BREAKDOWN VOLTAGE — The reverse voltage across a diode when the zener current is equal to a value specified by the manufacturer.

BREAKEOVER VOLTAGE (V_{BO}) — The forward bias at which a PNPN device suddenly switches from its forward blocking region into its forward high conduction region.

CONTROLLED RECTIFIER — A PNPN device designed to operate at relatively high power.

DIAC — A special breakdown diode which can conduct in both directions.

DYNAMIC IMPEDANCE (Z_z) — In connection with zener diodes, the opposition to changes of zener current. Also called DYNAMIC RESISTANCE.

DYNAMIC RESISTANCE (R_z) — See DYNAMIC IMPEDANCE.

FORWARD BLOCKING REGION — The forward biased state of a PNPN device in which the forward current is held to a low value by the high resistance of the middle junction.

FORWARD BREAKDOWN — In a PNPN device, the sudden change of junction resistance which occurs at the breakover voltage point, beyond which relatively low forward voltages produce relatively large forward currents.

GATE — The internal P-region of a controlled rectifier.

HIGH CONDUCTION REGION — The region in which a PNPN device conducts large forward current at low forward voltages after breakdown.

INTRINSIC STAND-OFF RATIO — The ratio of the interbase resistances of the unijunction. This ratio determines the internal voltage which appears at the emitter and controls the point at which conduction occurs.

NEGATIVE RESISTANCE — A characteristic whereby a decrease in voltage produces an increase in current, and vice versa.

PEAK POINT VOLTAGE (V_P) — The level of emitter voltage required to produce conduction in a UJT.

IMPORTANT DEFINITIONS (Continued)

SILICON CONTROLLED RECTIFIER (SCR) — A high-current controlled rectifier which uses silicon semiconductor material.

TEST CURRENT (I_{ZT}) — A typical operating zener current employed to specify the breakdown voltage.

TRIAC — A special controlled rectifier that can conduct in both directions.

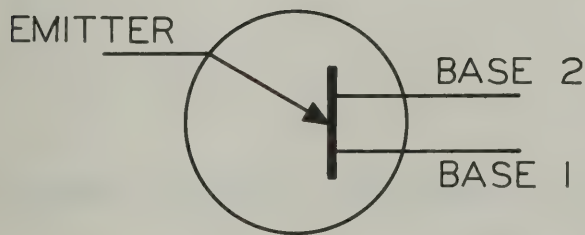
UNIUNCTION TRANSISTOR (UJT) — A special single PN junction device possessing negative resistance characteristics over a portion of its operating range.

ZENER CURRENT — Reverse current in the breakdown region.

ZENER DIODE — See BREAKDOWN DIODE.

PRACTICE EXERCISE SOLUTIONS

1. (b) reverse resistance is greater than the forward resistance.
2. (c) breakdown region.
3. True
4. (a) the rated breakdown voltage.
5. (b) diode D_1 .
6. (a) diode's resistance to decrease.
7. (c) the outer N-region is negative with respect to the outer P-region.
8. True
9. (b) holding current.
10. False — Increasing I_G decreases V_{BO} .
11. gate
- 12.



13. False — The unijunction transistor has one PN junction.
14. True
15. (a) forward biased.
16. (a) conducts.
17. The values of R_1 and C_1 .
18. (c) a positive pulse to the base.
19. True

PRACTICE EXERCISE SOLUTIONS (Continued)

20. Shifting the $I_G - V_{G2}$ phase relationship changes the length or time duration of the I_F pulses and therefore changes the amount of power supplied to the load.
21. (a) SCR_3 is on.
22. The amount of power can be controlled by varying the length of time the SCR conducts during a cycle of the supply voltage.
23. (b) only a portion of the positive alternation.
24. (b) minimum resistance.
25. C_1
26. (b) the full positive alternation .
27. (b) later during the positive supply alternations.
28. False — Maximum power is obtained with R_2 set to minimum, since the amount of phase shift is then minimum and the SCR fires early in the alternation. (See Figure 29D)
29. (b) R_3 .
30. True
31. (a) conducts.
32. The time required for C_1 to charge determines when UJT_1 conducts.
33. (b) controlled rectifier.
34. True
35. (a) more efficient.
36. (b) both alternations of the ac supply.
37. True
38. True

1103A

1. A - IN THE BREAKDOWN REGION OF A ZENER DIODE, A LARGE CHANGE OF CURRENT CAN OCCUR -- WITH A SMALL CHANGE OF VOLTAGE.

A zener diode exhibits a nonlinear resistance characteristic in the breakdown region. Small increases in voltage across the diode cause high conductivity, which implies a greatly lowered resistance. Since the resistance varies by a large amount in the breakdown region, the change in current through the diode is also inordinately large.

2. D - IN THE VOLTAGE REGULATOR CIRCUIT OF FIGURE 3, AN INCREASE IN V_{IN} CAUSES -- V_1 TO INCREASE.

Regulation is made possible by the voltage divider network consisting of R_1 and D_1 . Since D_1 is capable of responding in a nonlinear fashion to changes in V_{IN} , the greatest amount of voltage variations are dropped across R_1 .

3. A - CONDUCTION OCCURS IN A UNIJUNCTION TRANSISTOR WHEN -- THE EMITTER JUNCTION IS FORWARD BIASED.

Since there is only one PN junction in a UJT, this junction must be forward biased in order for conduction to occur, just as with any PN junction.

4. B - A PNPN DEVICE CHANGES FROM THE FORWARD BLOCKING REGION TO THE HIGH-CONDUCTION REGION AT A BIAS LEVEL CALLED THE -- BREAKOVER VOLTAGE.

The breakover voltage is the voltage which causes FORWARD BREAKDOWN of the reverse-biased junction in a PNPN device. When forward breakdown has been reached, the device operates in the high-conduction region.

5. A - AFTER THE BREAKOVER VOLTAGE HAS BEEN REACHED IN A PNPN UNIT, A LARGE FORWARD CURRENT CAN BE PRODUCED BY A RELATIVELY -- SMALL FORWARD VOLTAGE.

The large change in current associated with small changes in applied voltage is due to the nonlinear resistance characteristics of the high-conduction region.

6. B - A PNPN DEVICE IS SAID TO BE IN ITS "ON" STATE WHEN IT IS OPERATING IN ITS -- HIGH-CONDUCTION REGION.

For voltages below the breakover voltage, the only current through the device is a small leakage current across the reverse-biased junction. This current is negligible, and the device is effectively "off". The high-conduction region is the only area in a PNPN device where appreciable current flows.

7. D - INCREASING THE GATE CURRENT IN A FORWARD BIASED PNPN UNIT -- DECREASES THE BREAKOVER VOLTAGE.

The forward breakover voltage is dependent upon the gate current. Generally, a voltage which makes the gate positive with respect to the cathode of a PNPN device (such as an SCR) will produce gate current. If gate current is increased far enough, the breakover voltage is lowered to the anode-cathode voltage. This gate current is called the "gate current to fire".

8. D - IN FIGURE 28, C_1 CHARGES TO THE VALUE OF THE SOURCE VOLTAGE DURING THE NEGATIVE HALF-CYCLE. THE DISCHARGE PATH FOR C_1 DURING THE POSITIVE HALF CYCLE IS FROM THE CAPACITOR TO SOURCE TERMINAL NO. 1 THROUGH -- R_1 AND R_2 .

The charging path for C_1 on the negative half cycle is from source terminal No. 2 to C_1 , through D_1 and back to the source (terminal 1). The discharge path is from terminal No. 1 through R_1 and R_L , to the capacitor. Current cannot flow through D_1 during discharge because it is reverse biased during the positive half-cycle. Also, current does not flow through the load and SCR₁ until the capacitor discharges, thus making the SCR gate positive later in the cycle.

9. A - AN ADVANTAGE OF THE CONTROLLED RECTIFIER FOR SWITCHING IS -- THE LOW CONTROL CURRENT REQUIREMENTS.

The gate controls a relatively large amount of current through the device while itself requiring only small control currents.

10. D - A DEVICE DESIGNED TO DIRECTLY CONTROL LARGE AMOUNTS OF POWER IN AC CIRCUITS IS THE -- TRIAC.

A triac is a device which is the equivalent of two SCR's connected in parallel but with the polarities reversed. Thus conduction occurs through one of the SCR's on both the positive and negative half cycles.

PRACTICE EXERCISE SOLUTIONS (Continued)

20. Shifting the $I_G - V_{G2}$ phase relationship changes the length or time duration of the I_F pulses and therefore changes the amount of power supplied to the load.
21. (a) SCR_3 is on.
22. The amount of power can be controlled by varying the length of time the SCR conducts during a cycle of the supply voltage.
23. (b) only a portion of the positive alternation.
24. (b) minimum resistance.
25. C_1
26. (b) the full positive alternation .
27. (b) later during the positive supply alternations.
28. False — Maximum power is obtained with R_2 set to minimum, since the amount of phase shift is then minimum and the SCR fires early in the alternation. (See Figure 29D)
29. (b) R_3 .
30. True
31. (a) conducts.
32. The time required for C_1 to charge determines when UJT_1 conducts.
33. (b) controlled rectifier.
34. True
35. (a) more efficient.
36. (b) both alternations of the ac supply.
37. True
38. True

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1103A

Example: A well-known musical instrument is the
 A ☐ B ☐
 C ☒ D ☐
 A. fog horn. B. ear drum. C. trombone. D. door bell.

1. A ☒ B ☐ C ☐ D ☐ In the breakdown region of a zener diode, a large change of current can occur (A) with a small change of voltage. (B) only when forward bias is applied. (C) only if a large voltage change occurs. (D) because reverse bias produces forward current.

2. A ☐ B ☐ C ☐ D ☒ In the voltage regulator circuit of Figure 3, an increase in V_{IN} causes (A) an equal increase in V_O . (B) V_Z to decrease. (C) V_1 to decrease. (D) V_1 to increase.

3. A ☒ B ☐ C ☐ D ☐ Conduction occurs in a unijunction transistor when (A) the emitter junction is forward biased. (B) the emitter junction is reverse biased. (C) the emitter is connected to base 1. (D) base 2 is negative with respect to base 1.

4. A ☐ B ☒ C ☐ D ☐ A PNP device changes from the forward blocking region to the high-conduction region at a bias level called the (A) zener voltage. (B) breakover voltage. (C) reverse breakover voltage. (D) peak reverse voltage.

5. A ☒ B ☐ C ☐ D ☐ After the breakover voltage has been reached in a PNP unit, a large forward current can be produced by a relatively (A) small forward voltage. (B) large reverse bias. (C) small reverse bias. (D) low peak reverse voltage.

6. A ☐ B ☒ C ☐ D ☐ A PNP device is said to be in its "on" state when it is operating in its (A) reverse breakdown region. (B) high-conduction region. (C) forward blocking region. (D) reverse blocking region.

7. A ☐ B ☐ C ☐ D ☒ Increasing the gate current in a forward biased PNP unit (A) increases the PRV. (B) may drive the unit into reverse breakdown. (C) increases the breakover voltage. (D) decreases the breakover voltage.

8. A ☐ B ☐ C ☐ D ☒ In Figure 28, C_1 charges to the value of the source voltage during the negative half-cycle. The discharge path for C_1 during the positive half-cycle is from the capacitor to source terminal No. 1 through (A) D_1 , R_1 , and R_2 . (B) the load and SCR_1 . (C) SCR_1 , D_2 and D_1 . (D) R_1 and R_2 .

9. A ☒ B ☐ C ☐ D ☐ An advantage of the controlled rectifier for switching is (A) the low control current requirements. (B) its relatively large size. (C) the high voltage across it during conduction. (D) its very slow switching speed.

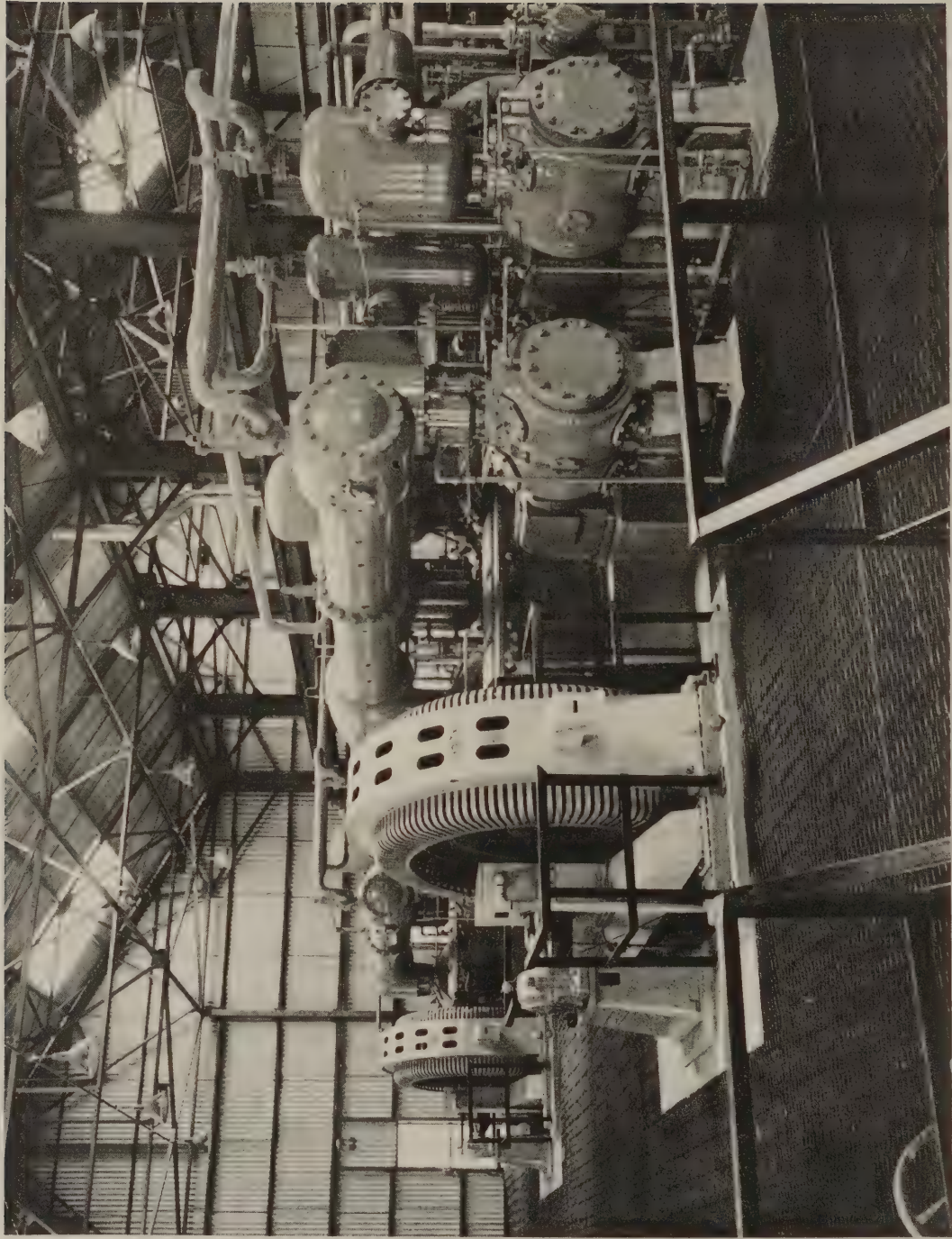
10. A ☐ B ☐ C ☐ D ☒ A device designed to directly control large amounts of power in ac circuits is the (A) zener diode. (B) two-layer device. (C) three-layer device. (D) triac.

MOTORS AND GENERATORS

1106

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





The wide variety of applications for motors and generators in industry is reflected in the almost infinite number of motor types and sizes in use. Shown here are two 1250 hp., 360 rpm, 4160 volt synchronous motors driving gas compressors.

Courtesy Electric Machinery Mfg. Co.

CONTENTS

Basic Generator Action	Page 5
AC Generators	Page 9
DC Generators	Page 11
Field Excitation	Page 13
Self Excited Generators	Page 14
Regulation	Page 16
Basic Construction	Page 17
Armature	Page 20
Motors	Page 22
The Motor Principle	Page 22
Torque	Page 24
DC Motors	Page 27
Series Motors	Page 28
Shunt Motors	Page 29
Compound Motors	Page 29
Motor Starting	Page 30
Motor Efficiency	Page 31
Three-Phase Power	Page 33
Synchronous AC Motors	Page 33
Three-Phase Induction Motors	Page 35
Single-Phase Induction Motors	Page 37
Shaded-Pole Motors	Page 39
Split-Phase Inductor Motors	Page 39
Split-Phase Capacitor Motors	Page 40

No man is free who is not master of himself.
—Epictetus

MOTORS AND GENERATORS

Motors and generators are electromechanical devices which perform important functions in electrical and electronic applications, and both may be called **TRANSDUCERS**. A transducer is a device which converts energy from one form to another. In the case of motors and generators, this energy is either electrical or mechanical. Generators are INPUT transducers since a generator converts nonelectric (mechanical) energy into electric energy. Motors are OUTPUT transducers since they convert electric energy into mechanical energy.

The mechanical energy supplied to the input of a generator is often from some type of **TURBINE**. A turbine is a type of rotary engine that provides rotational mechanical motion as a result of some type of fluid pressure. Figure 1 shows a simple turbine. The water pressure of falling water

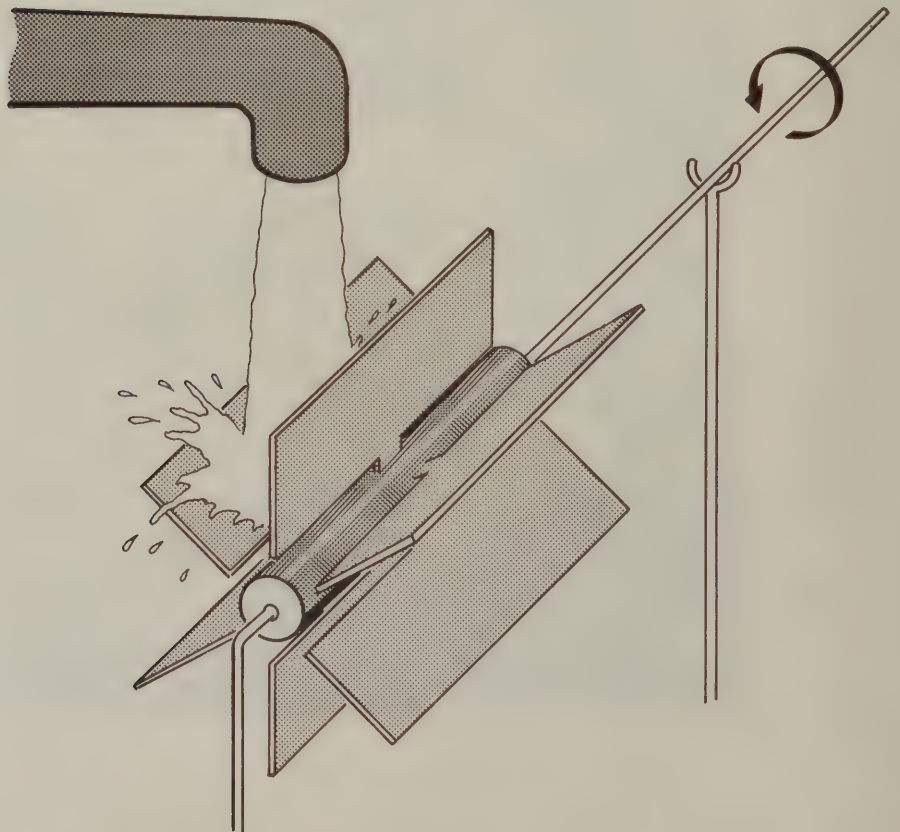


Figure 1

represents mechanical (kinetic) energy. The turbine blades transmit the energy to the shaft which rotates in the direction shown. Thus, a rotational mechanical force is produced which can be used as the input to a generator. This principle is used to produce electricity in large hydroelectric plants. The mechanical energy in these giant plants is supplied by natural water sources such as rivers, waterfalls and dams. The turbines drive gigantic generators which produce electricity for commercial and industrial use.

BASIC GENERATOR ACTION

GENERATORS produce voltages by **MAGNETOELECTRIC INDUCTION**.

These voltages, if applied to a complete circuit, will cause an electric current to flow. When a wire conductor moves through a magnetic field in such a way as to cut across the magnetic lines of force, a voltage is **INDUCED** in the wire. Voltage generated by induction is proportional to the number of flux lines cut per second. This voltage is the GREATEST when the conductor moves AT RIGHT ANGLES (90°) to the magnetic field. The voltage is ZERO when the conductor moves PARALLEL to the magnetic field.

The magnitude of the output voltage from a generator is determined by the following three factors:

1. The speed of the conductors through the magnetic field.
2. The strength of the magnetic field.
3. The number of conductors moving through the magnetic field.

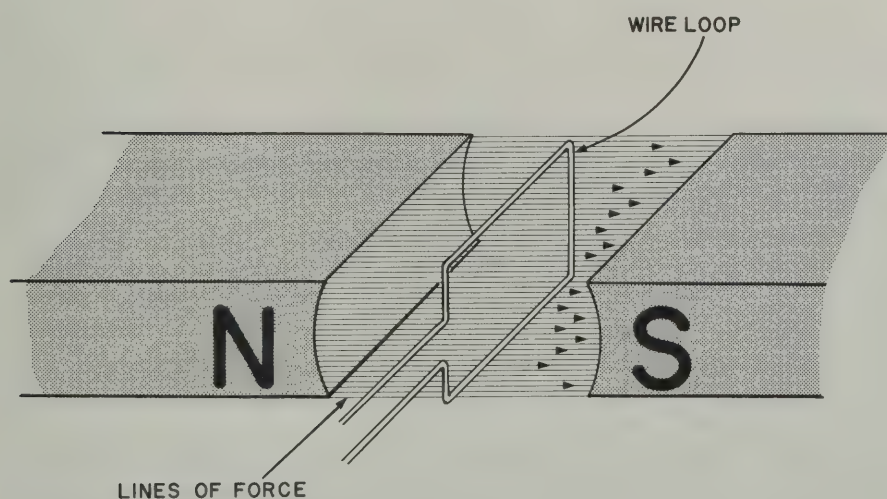


Figure 2

All three of these factors determine the number of flux lines cut per unit of time. It should be remembered that a voltage is generated by induction when the conductor and field move relative to each other. This means that either the field may remain stationary while the conductors move, or the conductors may remain stationary while the field moves. Figure 2 shows a single loop of wire placed in a magnetic field, and represents a simple **GENERATOR**. By connecting the generator to the turbine shaft shown in Figure 1, a basic **DYNAMO** is produced as shown in Figure 3.

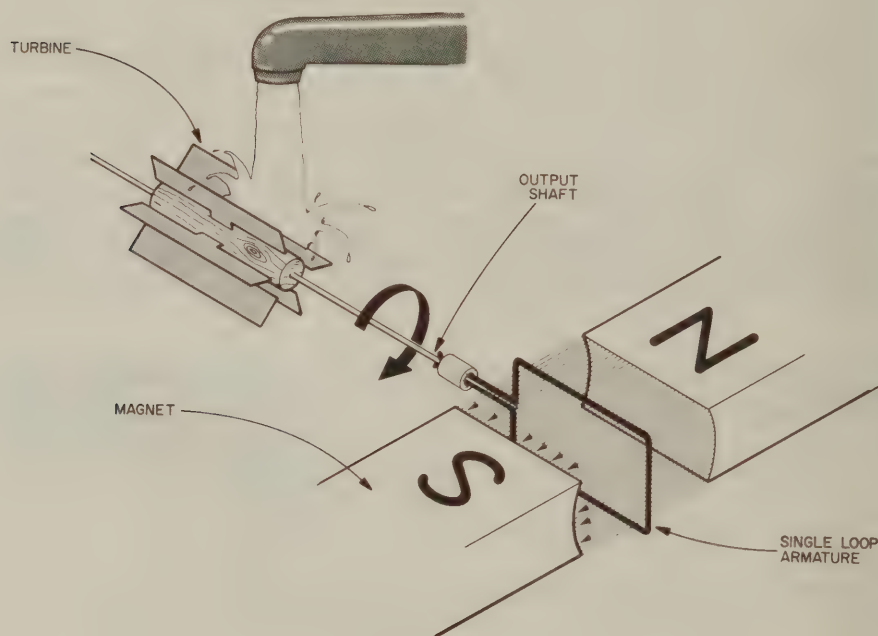


Figure 3

A dynamo is a turbine and a generator used together to convert mechanical energy to electric energy. The output voltage of any generator is taken from the **ARMATURE** windings. In the generator, the armature is attached to the rotating shaft. The rotating assembly including, in this case, the armature windings is called the **ROTOR**. The EMF induced in the armature windings can be connected to an external circuit which will permit a current to flow. Various arrangements used to supply the voltage to an external circuit will be discussed separately (ac-dc, shunt, series operation).

The armature output from any generator is ac. In order to understand the reasons for this, it is necessary to analyze the current flow characteristics of a conductor moving in a magnetic field. Current direction in a generator armature is determined by the **LEFT HAND GENERATOR RULE**, as

shown in Figure 4A. The first step in using the left hand generator rule is to determine the direction of the magnetic lines of force. Lines of force are considered to flow from the north pole to the south pole of a magnet. Since unlike poles attract, the lines of force flow from N to S between the two magnets as shown. If the N and S poles are not known, one can use a compass to determine the polarity. By placing the compass over the magnet, the S end of the compass will point to the N end of the magnet.

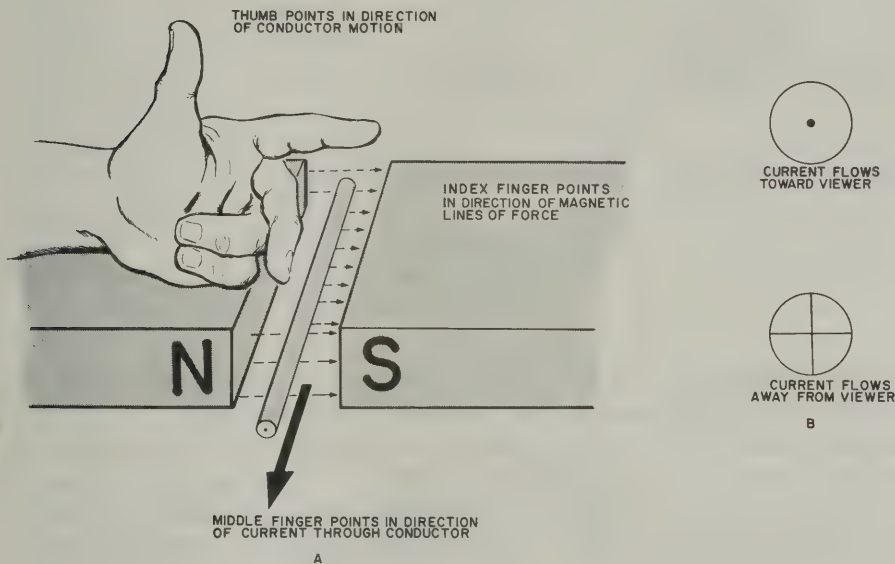


Figure 4

Next, assume that the conductor shown in the figure is initially below the magnetic field and is quickly moved up between the magnets. As the conductor moves up through the field, the conditions are as shown in the figure. Note that the thumb of the left hand points in the direction of conductor motion through the field; the index finger points in the direction of the magnetic lines of force (North to South); and the middle finger indicates the direction of current flow, as shown by the heavy arrow. The conductor is considered to be connected to a circuit, thus allowing current flow through it.

Figure 4B shows two symbols which are used to show the direction of current flow in a two-dimensional drawing. The circle is thought of as a conductor viewed from one end. A circle with a dot in the center indicates that the current in a conductor is flowing TOWARD the viewer. The dot is thought of as the tip or spear of an arrow coming toward the viewer, thus indicating that the current is flowing in the direction of the arrow. A circle with a "+" in it indicates that the current in a conductor is flowing AWAY FROM

the viewer. The + represents the stabilizer (feathered) end of an arrow traveling away from the viewer, thus indicating that current is traveling in that direction.

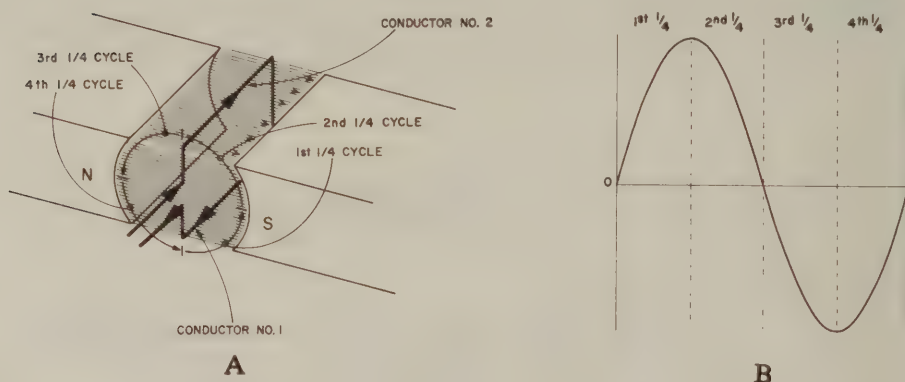


Figure 5

Figure 5A illustrates an armature coil rotating in a magnetic field. Figure 5B shows the output obtained for one revolution of the armature. As the coil rotates, it cuts more or fewer lines of force depending upon its position at any given instant of time. With the armature rotating through the vertical position as shown in Figure 5A, conductors 1 and 2 are moving parallel to the magnetic field. Since no flux lines are being crossed or cut, the armature output under the condition is zero, as indicated by the beginning of the trace in Figure 5B. As the coil continues to rotate through the first $\frac{1}{4}$ cycle, the conductors gradually cut across more and more flux lines, and current flows in the loop of wire as shown by the HEAVY arrows in Figure 5A. This can be determined by using the left hand generator rule since side No. 1 of the loop moves UP and side No. 2 moves DOWN. At the end of the first $\frac{1}{4}$ cycle, the conductors are moving at right angles to the flux lines, and are then cutting the maximum number of lines per second. The armature output is maximum at the end of the first $\frac{1}{4}$ cycle as shown by the trace in Figure 5B. During the second $\frac{1}{4}$ cycle, the conductors gradually cut across fewer and fewer flux lines until the conductors once again are moving parallel to the field. The armature output returns to zero after 180° of rotation as shown in Figure 5B.

Note that the loop of wire is again oriented in the vertical position between the magnets. Conductor No. 2 is now on the bottom and conductor No. 1 is on top. Therefore, to find the direction of current in the loop all we have to do is change the labeling of the conductors in Figure 5A. Current still flows in the same direction through the upper wire. Just remember that conductor No. 1 is now the upper wire. Thus, as rotation continues through the third and fourth quarter-cycles, THE CURRENT CHANGES DIREC-

TION. This is also shown in Figure 5B by the negative swing of the trace. By using the left hand generator rule, you will see that the current is again represented by the HEAVY arrows. Thus, the current changes direction each half-revolution, or half-cycle, and an alternating current is produced.

AC GENERATORS

In order to take the ac output from a generator, the ends of the coil must be attached to **SLIP RINGS** which rotate along with the armature. A simple ac generator is shown in Figure 6. Stationary conductors called **BRUSHES** are placed in contact with the slip rings and allowed to slide along the surface as the rings rotate. In this way, brushes and slip rings provide a connection between the armature coil and an external load which uses the power supplied by the generator. Brushes usually contain carbon or graphite to maintain low friction without cutting into the metal slip rings. As the brushes ride on the slip rings, they slowly wear away and must be replaced periodically. Since the brushes wear to the shape of the slip rings, the resistance between them and the slip rings can be very low.

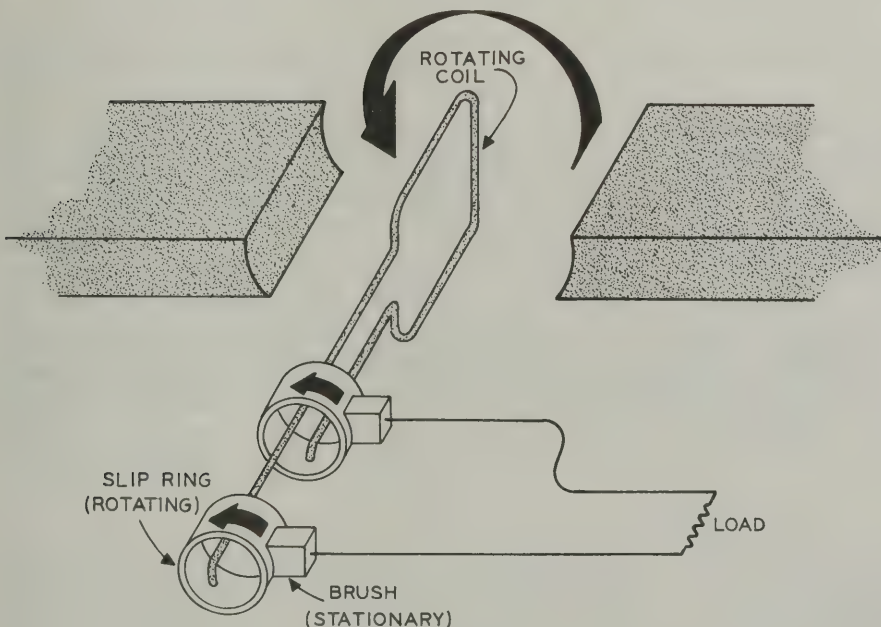


Figure 6

AC generators are also called **ALTERNATORS**. Alternators provide most of the electric power used in homes and industry. These machines are very

large and generate large amounts of power. As pointed out previously, a voltage may be induced in a conductor regardless of whether the conductor moves or the field moves. The simple generator in Figure 5 uses a stationary field called the **STATOR**, and a rotating armature, or rotor. The output voltage is taken from the rotor. However, with larger machines where large amounts of power are generated, it is impractical to take the output power through the slip rings and brushes. The slip rings and brushes are better suited to low power due to the special nature of the materials used. We know that a small magnetic field will produce a large voltage if a large number of conductors cut the lines of force, or if the lines of force are cut at great speed. Therefore, many large alternators reverse the roles of the stator and rotor. In these machines, a small dc voltage is applied through the brushes and slip rings to a rotating **ELECTROMAGNET**. The dc voltage is called **DC EXCITATION**. The electromagnet produces the field, and, instead of turning the armature winding, the rotor turns the **FIELD WINDING**. The field coil is wound on a soft iron core and, with a steady dc voltage applied to the slip rings, a steady electromagnetic field is produced. The **ARMATURE** coils are then wound on the **STATOR**. The rotor can be driven at high speed, or a large number of armature windings can be used, and the high voltage can be conveniently taken from the stationary armature windings. Since relative motion is produced between the rotating field and the stationary armature in the same way as for the reverse arrangement, the output is ac just as in the previous example.

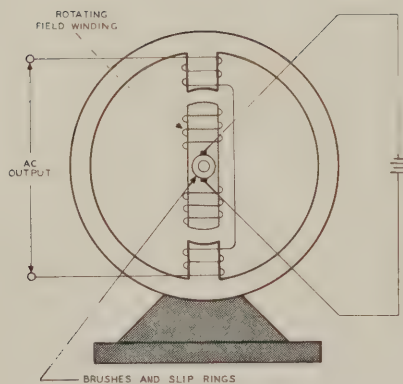


Figure 7

Generators which operate on the principles of electromagnetic induction operate with a steady field which is normally provided by dc excitation. The alternator shown in Figure 7 is typical of a 2 pole **SYNCHRONOUS ALTERNATOR** which uses dc excitation from an external source. The shaft is driven by a mechanical force at a constant speed of 1,800 or 3,600 rpm. A synchronous alternator is the type used to produce electricity for ordinary household purposes.

The output ac frequency of an alternator is determined by the speed of the rotor and the number of poles in the alternator. The formula for determining the ac frequency is

$$f = \frac{p \times \text{rpm}}{120}$$

where

f = the frequency in hertz

p = the number of magnetic poles in the alternator

rpm = the rotor speed of the alternator in rotations per minute.

As an example, consider a four-pole synchronous alternator whose rotor is being driven at a constant 1,800 rpm. The output frequency is

$$f = \frac{4 \times 1,800}{120} = 60 \text{ Hz}$$

DC GENERATORS

In a dc generator, the ac output from a rotating armature is converted to pulsating dc by using a **COMMUTATOR** in place of slip rings. Figure 8 shows a simple dc generator. The rotating coil and the magnetic field are the same for both ac and dc generators. The difference is in the method used for removing the voltage induced in the armature. In an ac generator the armature coil is attached to slip rings which make contact with brushes. In a dc generator the armature coil is attached to a commutator which also makes contact with brushes.

A commutator is basically a slip ring which is split into two (or more) parts. These parts are called commutator **SEGMENTS**. The segments are insulated from each other and from the shaft. Each end of the armature coil is attached to one of the segments, as shown in Figure 8A. The brushes make contact to opposite sides of the commutator segments.

Notice also that the brushes "bridge" the commutator segments when the coil is in the position shown. This short-circuits the coil for an instant as the coil rotates. However, since the coil conductors 1 and 2 are moving parallel to the magnetic field at this point, no voltage is induced in the loop. The area in a generator where no emf is being induced is called the **COMMUTATING** or **NEUTRAL PLANE**. The brushes are always set so that they touch two commutator segments while passing through the neutral plane. If the brushes were to bridge the commutator segments while an emf was being induced in the coil, a heavy current would flow in the coil. Since the coil and brushes have a very low resistance, the armature coil would be damaged unless commutation takes place in the neutral plane.

Pulsating dc is obtained from a dc generator because the connections from the armature coil are reversed each half-cycle. This can be shown by using the left-hand generator rule in conjunction with Figure 8A. With the armature coil in the neutral plane, coil rotation produces no armature output as shown at the beginning of the trace in Figure 8B.

As the armature rotates past the neutral plane, the coil conductors begin to cut across the flux lines, and current begins to flow in the armature

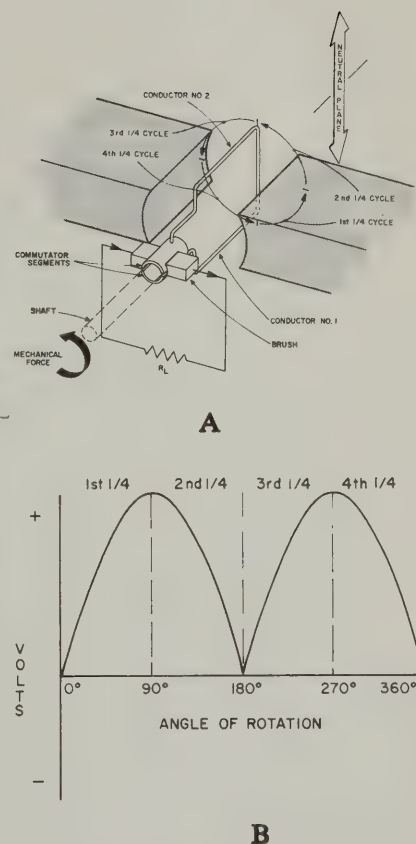


Figure 8

circuit. The current increases steadily until the armature coil is moving directly across the face of the magnets. At this time the coil conductors are cutting the flux lines at the maximum rate. The armature output is represented at this point by the trace in Figure 8B at the end of the first $\frac{1}{4}$ cycle (90°). During the second quarter-cycle, the armature coil current decreases steadily until the armature is once again oriented in the vertical position. When reaching the end of the second quarter-cycle, the coil has rotated 180° and the conductors are again moving parallel to the flux lines. Armature output returns to zero as shown by the trace at 180° in Figure 8B.

Just as in an ac generator, the armature coil current **REVERSES DIRECTION** at the beginning of the third quarter cycle. (Coil conductor No. 1 moved **UP** through the magnetic field during the first and second quarter cycles, and moves **DOWN** through the field during the third and fourth quarter cycles. Using the left-hand generator rule, the thumb must be pointed down for conductor No. 1 during the third and fourth cycles. Since the direction of the field remains the same, the current through conductor No. 1 must reverse, as shown by the middle finger of the left hand.)

Up to this point, the operation for dc generators has been the same as for ac generators. However, note that just as the coil current reverses, the commutator reverses the electrical connection to the load resistor R_L . This maintains current flowing in the load circuits in the same direction as shown by the arrows in Figure 8A.

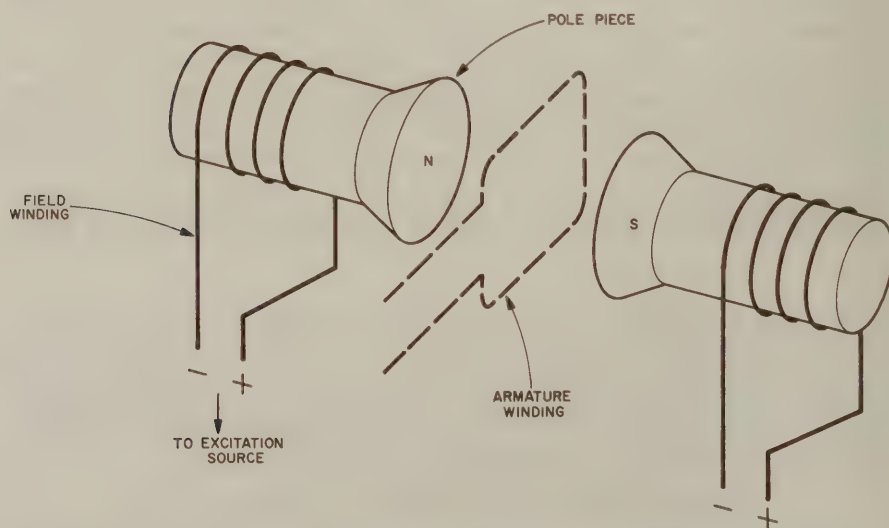


Figure 9

Field Excitation

The permanent magnet, shown in the previous examples, provides the magnetic flux required to induce voltage in a moving conductor. However, an electromagnet may be used to supply the magnetic field in a generator. These electromagnet coils are called **FIELD WINDINGS**. The field windings are shown in Figure 9. The field coils are wound on **POLE PIECES** which act as the core of the electromagnet. Current through the coils produces an electromagnetic flux which uses the core, or pole piece, as a return path for a complete electromagnet circuit. The magnetic flux from the field cuts the armature, inducing a voltage in it. Thus, the pole piece becomes temporarily magnetized, and acts as a permanent magnet in a generator.

Unlike the permanent magnets in the generators of Figures 6 and 8, field windings require a current source in order to build up the field. This current is called **FIELD EXCITATION** current. When the excitation current is supplied from a separate source of dc such as another generator or a battery, the generator is **SEPARATELY EXCITED**. A separately excited generator is shown in Figure 10.

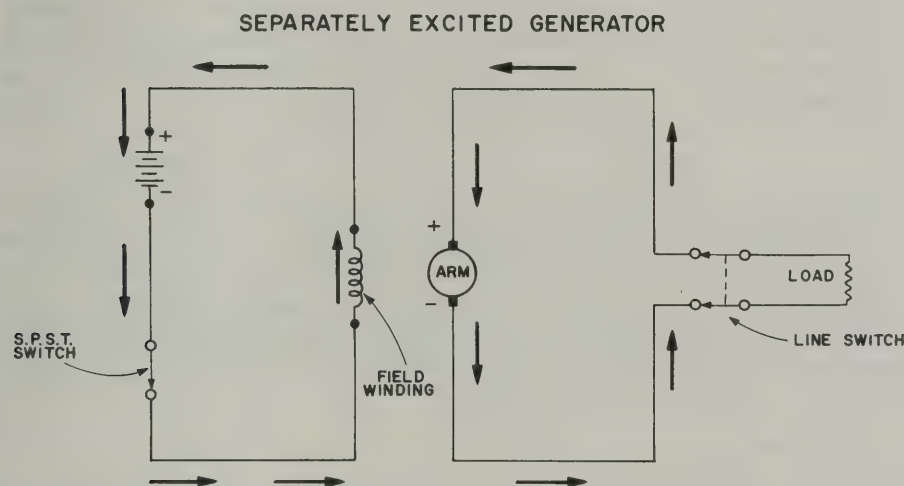


Figure 10

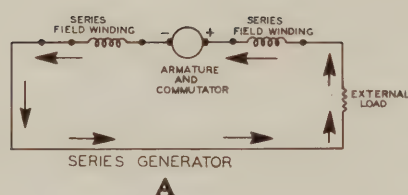
The excitation circuit is independent of the armature and load circuit. The strength of the separate excitation field depends upon the amount of current supplied to it.

Self Excited Generators

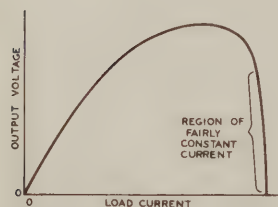
The current required for the electromagnetic field of a dc generator may be taken from the armature of the generator itself. This is called **SELF-EXCITATION**.

The self-excited generator uses part of the voltage developed in the armature to produce current in the field windings. This is possible once the generator is running because the armature, being turned mechanically, produces a voltage which forces current through the field coils to produce an electromagnetic field. The lines of force in the electromagnetic field are then cut by the rotating armature which supplies more current to the field coils.

A natural question which arises is, "How does a self-excited generator start generating voltage?" The reason that the generator is self-starting is due to **RESIDUAL MAGNETISM**. The pole pieces, acting as cores of an electromagnet, retain a small amount of magnetism even after the excitation current is removed. Thus, after the generator has been in use for a period of time, the pole pieces become magnetized. This magnetic field, although weak, is sufficient to induce a small voltage in the rotating armature as the generator is started up. The armature then supplies a small current to the field coils. This process continues with the current and electromagnetic field increasing until the generator reaches its maximum output rating.



A



B

Figure 11

Another method for starting a generator is to use **EXTERNAL EXCITATION** during the first few revolutions of the armature. The external excitation current is removed when the self-excitation current becomes sufficient to maintain generator output. This technique is called **FLASHING THE FIELD**. External excitation may be required when the pole pieces have lost their residual magnetism and the self-excited generator cannot supply starting current to its field windings.

Two types of self-excited generators are the **SERIES** and **SHUNT** generators. A series generator is one in which the armature windings, field windings and the load are connected in series, as shown in Figure 11A. Since all the current produced by the generator flows through the field windings and the load, any variation in the load resistance will cause the generator output to vary.

The voltage output of a self-excited series generator continues to rise (after being started by residual magnetism or by flashing the field) until the pole pieces become **SATURATED**. This means that the iron in the pole pieces cannot hold more lines of force, and any additional current through the

The following Practice Exercise questions cover the subjects which you have just studied. They are:

BASIC GENERATOR ACTION

AC GENERATORS

DC GENERATORS

FIELD EXCITATION

1. A device which converts energy from one form to another is called a _____.
2. A machine which uses a turbine and a generator to convert mechanical energy to electric energy is called a _____.
3. The rotating armature coil of Figure 5 cuts across the maximum number of lines of force twice each revolution. True or False?
4. The only time that a voltage is induced in a rotating coil is when it moves at right angles to the magnetic field. True or False?
5. What is the output frequency of a two-pole synchronous alternator when the rotor speed is 3600 rpm?
6. In a dc generator, the voltage induced in the armature is ac, but the commutator converts the output to dc. True or False?
7. The brushes are allowed to "bridge" two commutator segments only in the _____.
8. A commutator reverses the direction of current flowing in the armature in a dc generator. True or False?

field windings will not produce more flux. The output voltage of the generator is MAXIMUM at saturation.

Figure 11B shows a graph of the output voltage characteristics of a series generator. For increasing values of current required by the load BEFORE SATURATION, the output voltage increases. This is because the load and the field windings are in series. When the field windings reach saturation, however, increased current requirements by the load cause the voltage to drop off sharply. Thus we see that changing load conditions drastically affect the output voltage of a series generator, causing it to increase or decrease depending upon the operating point. For this reason, it should be remembered that series generators are used in CONSTANT CURRENT applications. That is, the CURRENT requirements of the load should be unchanging, and therefore as constant as possible.

The operating characteristics of a shunt dc generator are different from a series generator, as shown in Figure 12. Since the load and the field windings are in parallel and not in series, changing load current affects the armature voltage differently. In a shunt generator, increased current drawn through the load always reduces the generator output voltage. This is true because as more current flows through the load circuit the voltage drop across the armature increases, the output voltage decreases, and less current flows through the field windings. With less current flowing through the field windings, fewer electromagnetic lines of force are produced. This, in turn, reduces the voltage induced in the armature. Since the armature and the load are in parallel, the voltage across the load drops as more current is drawn. This process continues as the load increases, and if the load current continues beyond a critical value (rated current), the output voltage drops almost to zero as shown in Figure 12B. For this reason, a shunt dc generator should be used in CONSTANT-VOLTAGE APPLICATIONS. That is, the VOLTAGE requirements of the load should be unchanging, and as constant as possible. Overloading a dc shunt generator can burn out the armature, causing generator failure.

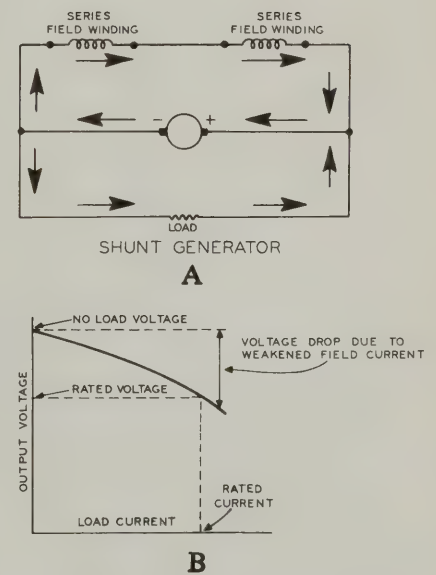


Figure
12

Many applications require generators with features of both the series and shunt types. As a result, the compound generator evolved. These generators use both series and shunt fields. Varying the strength of these fields allows for a wide variation in the output characteristics of the generators.

A FLAT-COMPOUNDED generator is one whose output voltage remains relatively constant even though the generator load changes. In an OVER-COMPOUNDED generator the series field predominates and the output voltage rises somewhat as the load current increases. If it is desired to have the generator voltage decrease slightly with increases of load current, the effects of the shunt winding are more prominent than those due to the series

winding. This latter machine is called an UNDER-COMPOUNDED generator.

Regulation

Since the voltage output of a shunt generator is critical, it is desirable to use some form of automatic voltage regulation. A voltage regulator is a device which senses a change in output voltage, and then restores the voltage to its original value. Figure 13A shows one method for regulating generator voltage. The regulator includes a variable resistor which is inserted in series with the field windings. Another set of leads from the regulator is connected across the load. When the voltage across the load changes, the change is sensed by a circuit in the regulator. This can be done by using a balanced bridge circuit, or some type of semiconductor or tube control device. The regulator then varies the series resistance directly as the load voltage changes. That is, if the load voltage decreases, the series resistance decreases. This permits more current to flow through the field windings, thus restoring the output voltage to its original value. Figure 13B shows a different method

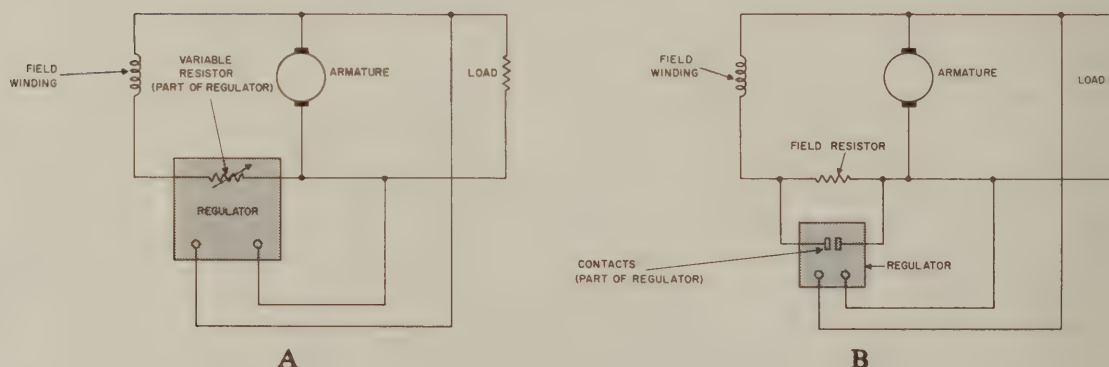


Figure 13

for regulating the output of a generator. This system requires a voltage sensing device across the load as before. However, the regulator places a fixed-value resistor in series with the field windings. By using a set of vibrating contacts across the resistor, the AVERAGE VALUE of current through the field winding is controlled. When the contacts close, the resistor is shorted out and field current increases. The contacts vibrate in such a way as to vary the amount of time they are closed compared to the amount of time they are open, depending upon whether the output voltage goes up or down. To illustrate this, suppose the output voltage suddenly increases due to a change in load current. The voltage sensing device senses this change and causes the vibrating contacts to open (not touch) for longer periods of time, causing the resistor to be in the circuit longer, thus reducing

field current. This reduces the current through the field windings, and restores the output voltage to the original value.

Basic Construction

The simple one-loop generator used in the examples thus far would be inefficient and impractical as an actual machine. The physical construction of a generator determines to a large extent the operating characteristics. Figure 14 shows the major parts of a dc generator. These are the frame, armature assembly, commutator assembly, brush assembly, pole pieces, field coils, shaft, and end bells. The frame not only supports the generator but provides a path for conducting magnetic lines of force. By using a high-

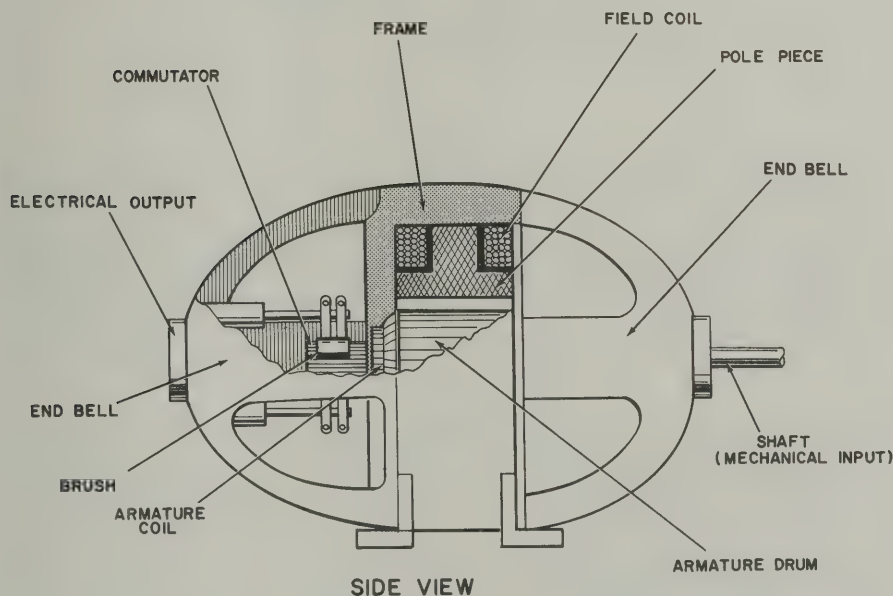


Figure 14A

permeability material in the frame, the electromagnetic field between the pole pieces is strengthened. An open-frame generator consists mainly of a circular frame with the essential induction and commutator components inside. A semiclosed-frame generator is essentially an open frame with screens over the ends. A closed-frame generator has end bells attached to the frame and provides additional protection for the moving parts inside. The drawings of a complete generator in Figure 14A is a closed-frame generator.

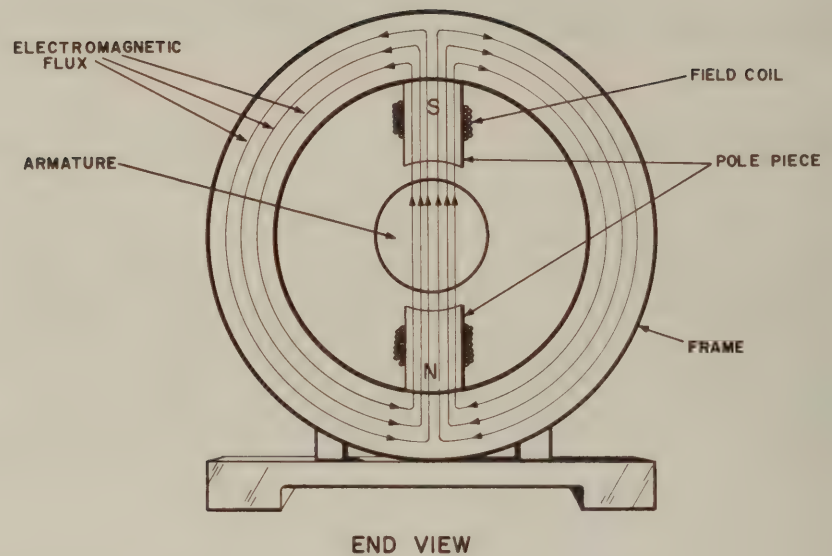
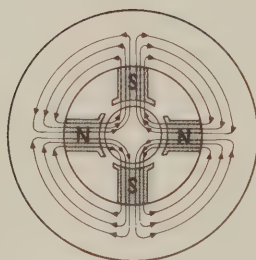
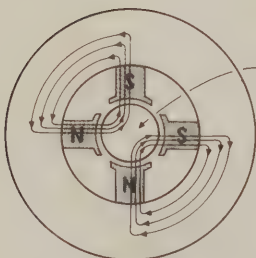


Figure 14B

The end view of Figure 14B shows the generator with the end bells removed. The pole pieces are shown attached to the frame supporting the field windings. This represents a two pole generator with the poles located opposite each other. The poles are placed directly opposite each other because, for a two pole generator, the armature can then cut the maximum amount of flux.



CORRECT
A



INCORRECT
B

Figure 15

Generator poles are used in pairs. Most generators use more than a single pair of poles. Figure 15 shows a four-pole generator and the arrangement of its pole pairs. Additional poles strengthen the magnetic field and provide more lines of force in the armature. Note also that the frame provides a good magnetic path return for the field windings. Another difference between double- and multi-pole generators is the pole placement. In a two-pair generator, N and S poles are not placed opposite each other. Instead, N and S are placed NEXT to each other. This provides electromagnetic flux throughout the frame and equally distributed within the armature. Figure 15 illustrates the difference between correct and incorrect pole placement. The incorrect placement produces flux gaps, and no voltage is generated when the armature passes through these areas.

To show another advantage of using additional poles in a generator, we go back to our simple generator model. Figure 16A shows two single-loop coils placed at right angles to each other. The commutator is split into four segments, and one end of each coil is attached to a commutator segment. With this arrangement, an emf is induced in one of the coils at all times. As

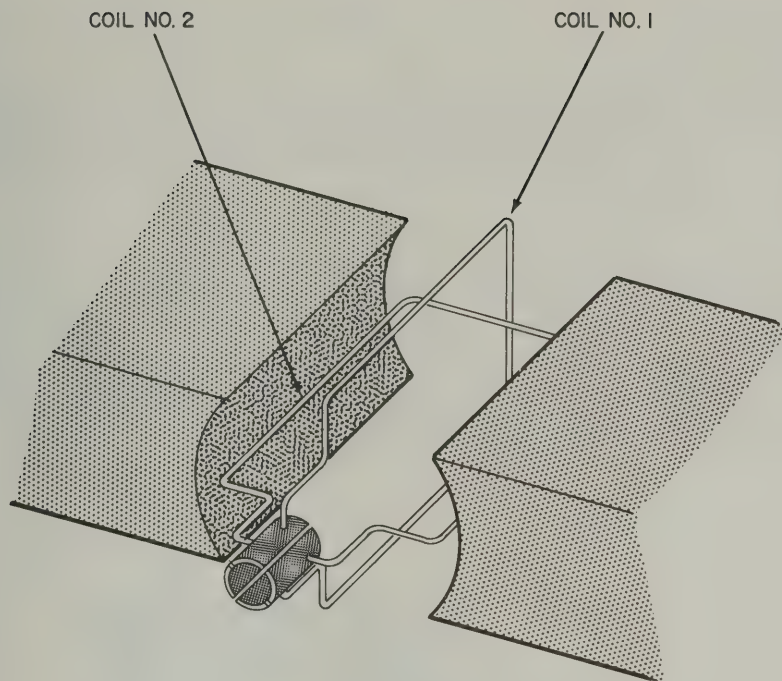


Figure 16A

coil No. 1 is the commutating plane, maximum emf is induced in coil No. 2. Although the output is still pulsating dc, the output is smoother than with a two-pole generator. This is shown in Figure 16B. Most generators use many poles rather than two or four, which produces a still smoother output.

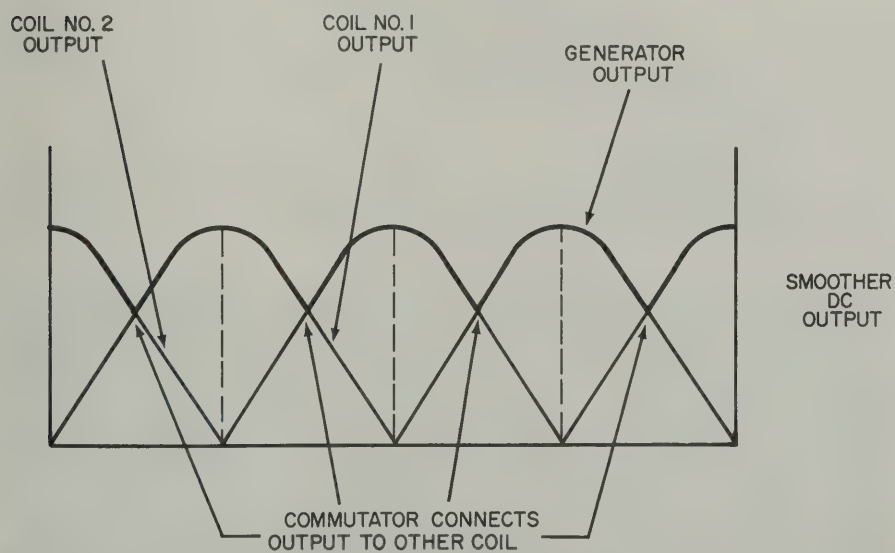


Figure 16B

Armature

The armature is basically made up of coils of wires. It is the method in which they are constructed that gives a generator armature its efficient operation. Most modern generators use a drum armature, as shown in Figure 17. The individual coils are wound on a core which resembles a drum. Instead of being wound on the surface, however, the coils are recessed in slots just a fraction of an inch below the outer drum surface. This protects the coils and reduces the chance that a coil will touch one of the pole pieces as it is rotating. Since the drum can be precision machined, it rotates very close to the pole pieces without actually touching them. This also reduces the air gap between the coils and the pole pieces, which permits a greater emf to be induced in the coils.

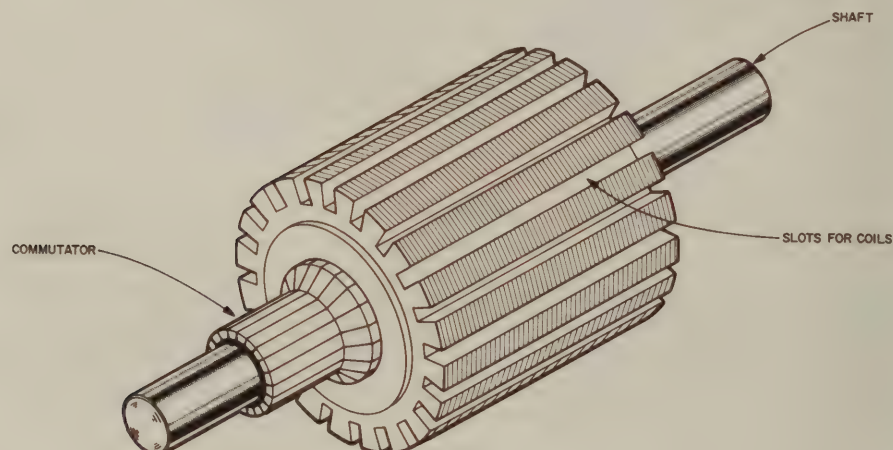


Figure 17

Thus, the drum actually serves two purposes. One is to support the coils. The other is to increase the efficiency of the armature.

The armature coils are wound on the armature core in such a way that the two slots (required for each coil) will be the same distance apart as adjacent magnetic stator poles. The distance between poles is called the **POLE PITCH** or **POLE SPAN**. Since adjacent magnetic poles are of the opposite polarity, the emf induced in one side of the coil is reinforced by the emf induced in the other side of the coil.

When the distance between the two sides of an armature coil is not equal to the distance between adjacent magnetic poles, the winding is called a

FRACTIONAL PITCH WINDING. Fractional pitch windings are used when it is necessary to conserve copper. The emf induced in a fractional pitch winding is less than that induced in a winding with a pitch equal to the pole pitch. This is because the voltages induced in the two coil sides do not reach their maximum values at the same time.

There are two ways to connect the coils together on an armature—**LAP WINDING** and **WAVE WINDING**. We will use **SIMPLEX LAP WINDING** to illustrate the first type. Simplex lap winding, shown in Figure 18A, uses an arrangement when the ends of each coil are connected to adjacent commutator segments. The effect of this is that all the coils are connected in series, and the armature acts as one very large coil, as shown in Figure 14. In a simplex lap winding, a brush shorts out the two ends of a single coil.

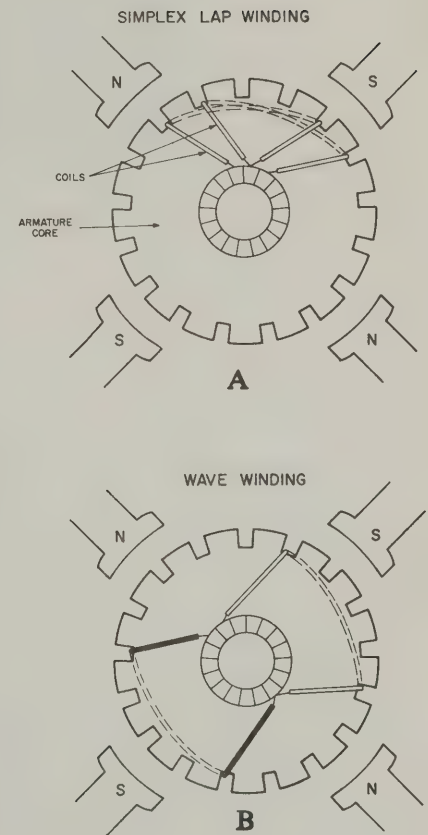
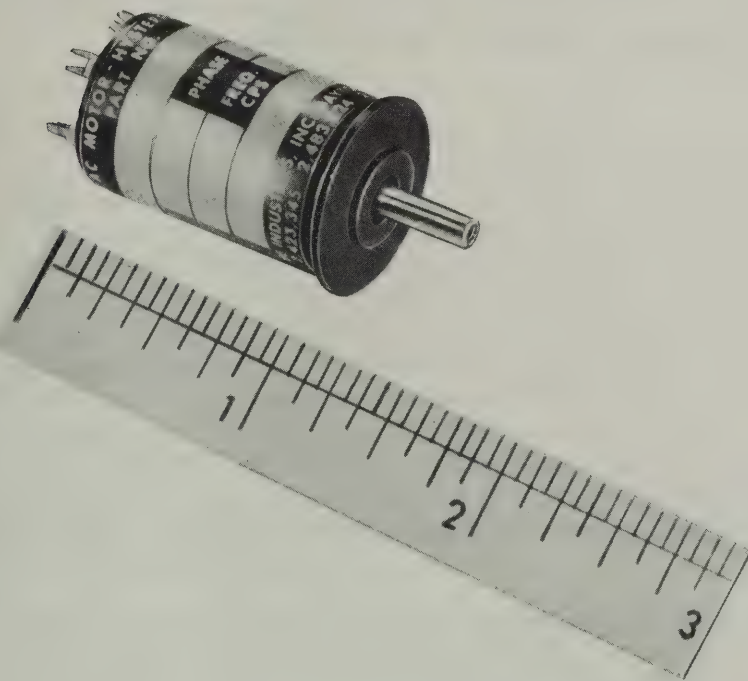


Figure
18



A subminiature synchronous AC motor measuring only 5/8" in diameter by 1-3/16" long.

Courtesy Globe Industries, Inc.

In wave winding, shown in Figure 18B, the ends of each coil are connected to commutator segments two pole spans apart. Instead of shorting a single coil, a brush shorts a group of coils in series. The number of coils depends upon the number of armature poles. The number of coils a brush will short is equal to the number of pairs of magnetic poles.

MOTORS

Motor construction is physically similar to generator construction. In fact, a basic dc generator can also function as a dc motor. Both dc motors and generators use external magnetic field excitation, armature coil windings, and commutators.

The Motor Principle

Motors operate because of the interaction between magnetic and electromagnetic fields. Figure 19A shows an arrangement which can be used to analyze these forces. First, however, we will briefly review the **LEFT-HAND CURRENT RULE**. The left-hand current rule states that current flowing through a conductor in the direction indicated by the thumb of the left hand produces electromagnetic lines of force around the conductor in the direction indicated by the curled fingers of the left hand.

Another rule is also used in determining the forces acting upon a motor. This rule is called the **RIGHT-HAND RULE FOR MOTORS**. The right-hand rule for motors is similar to the left-hand generator rule. That is, the fingers of the RIGHT hand are used to represent motor forces in the same way that the left hand represents generator forces. The difference is (besides the use of the opposite hand) that the generator rule is used to find CURRENT direction, while the motor rule is used to find the direction of MOTION. Remember that in a generator, the direction of armature (coil) motion was predetermined by a mechanical turning force. The direction of current flow in the armature, however, was unknown. In a motor, however, the situation is just opposite that of a generator. The direction of current through a motor armature is predetermined by an external current source. It is the direction of armature rotation which is unknown in motor operation.

To determine the direction of armature rotation, the thumb, index, and middle fingers of the right hand are held in a configuration similar to that shown in Figure 4A.

Since motor operation is dependent upon the interacting forces of magnetic and electromagnetic fields, we must use TWO rules in order to determine the DIRECTION of these fields. One is, of course, the right-hand rule for motors, and the other is the left-hand current rule. Referring again to Figure 19, use the left-hand current rule to establish the direction of the field encircling the conductor. Also remember that the flux travels from N to S in the stationary field. Current direction through the conductor in Figure 19 is constant in the direction of the arrow marked "current flow". The direction of the field around the conductor is indicated by the curled

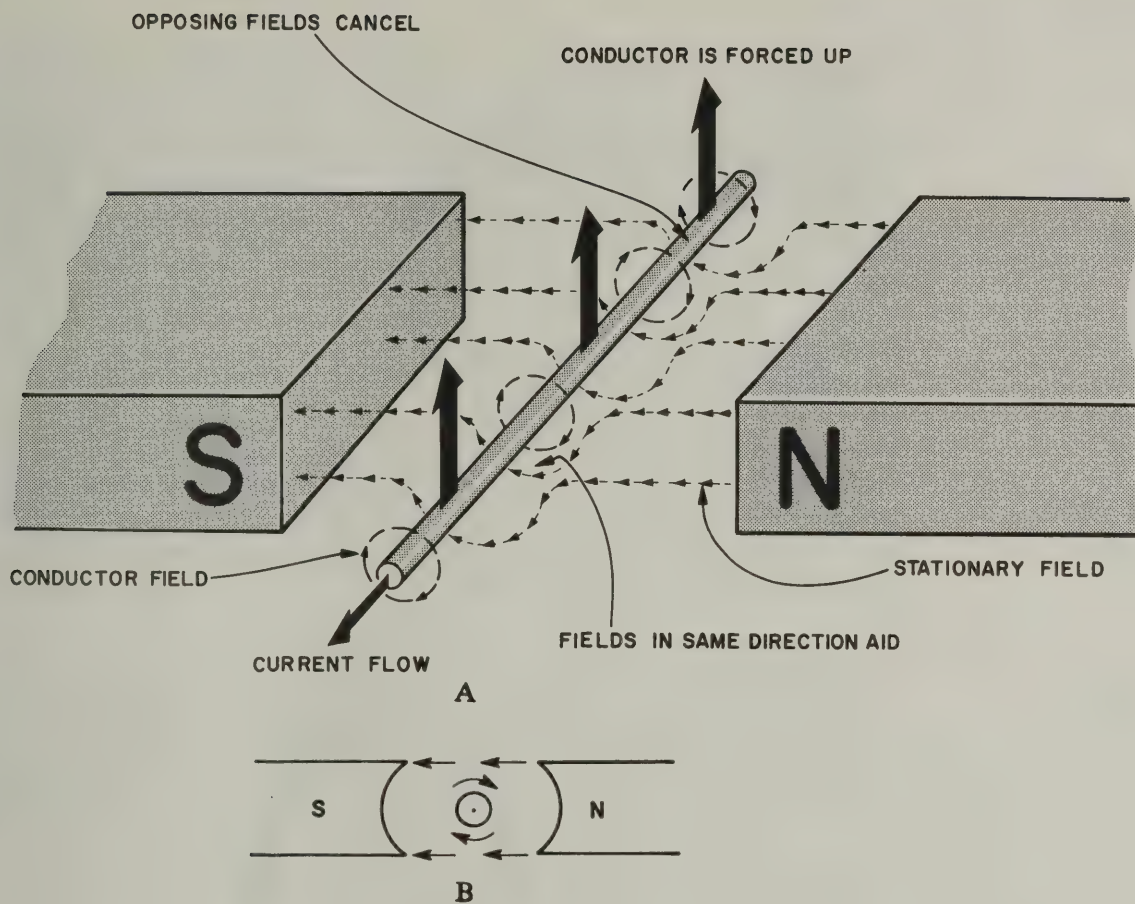


Figure 19

fingers of the left hand. The direction of the stationary field is indicated by the index finger of the right hand. Thus, using the two rules, current flow direction is indicated by the **THUMB** of the left hand rule, and the **MIDDLE FINGER** of the right hand rule.

Note that the conductor appears to be lying in a trough, or dent, in the magnetic field. This is caused by the CANCELLING effect of the magnetic and electromagnetic fields at the point where the fields oppose each other. In this case, the cancellation point is at the top of the conductor and the net result is a weakened field at this point. This is shown in Figure 19B. At the bottom of the conductor, however, the fields are in the same direction and **AID** each other. Therefore they are reinforced, becoming much stronger. The effect of this is an unbalance of forces on the conductor, and the conductor moves upwards as shown by the three heavy arrows. Note that this is the direction indicated by the **THUMB** of the right hand. This proves that the thumb of the right hand can be used to determine the direction of motion in motor action.

Torque

The force which is produced by the motor action just discussed is called **TORQUE** (pronounced "Tork"). Torque is a turning force. From the simple explanation given for illustrating motor action, we can determine that a turning force is produced when the conductor of Figure 19 is formed into a loop. Note that in Figure 20, the loop, or coil, forms two conductors in series, which are labeled conductor No. 1 and conductor No. 2. For conductor No. 1, the conditions are the same as for that of the straight conductor in Figure 19, and the conductor is forced up. For conductor No. 2, the current through the conductor is the reverse of that in Figure 19B. This means that the magnetic and electromagnetic fields cancel at the bottom of the conductor, and aid at the top. This situation is the reverse of the forces acting upon conductor No. 1, and this side of the loop is forced down. This upward pressure on conductor No. 1 and downward pressure on conductor No. 2 makes the coil turn as shown.

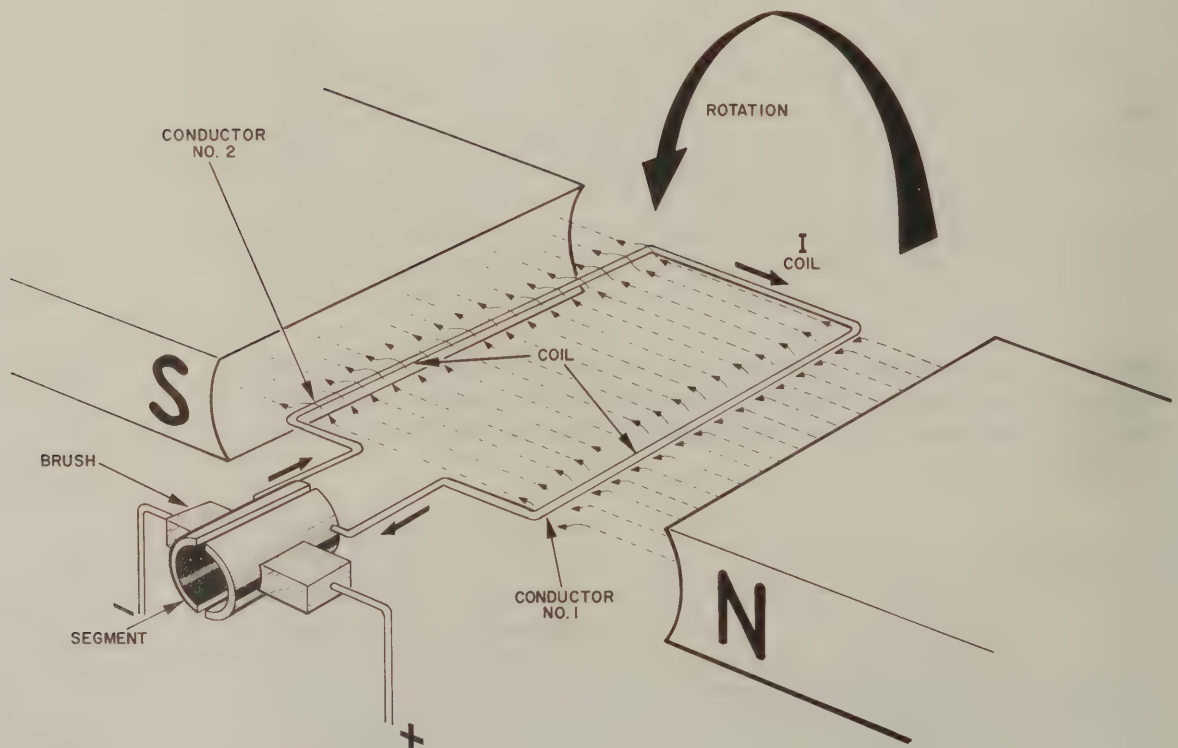


Figure 20

The amount of torque generated in this rotating motion depends upon the **STRENGTH OF THE MAGNETIC FIELD**, and the **DISTANCE OF THE**

SIDES OF THE COIL (conductors No. 1 and 2) FROM THE CENTER OF ROTATION. Also remember that the strength of any electromagnetic field depends upon the number of turns in the coil and the current through the coil. To illustrate a simple torque calculation, consider Figure 21. Shown is a board which is just balanced at the center. When a 10 lb. weight is placed two feet from the center of the board, a turning force is produced about the center, or pivot point. The torque is found by calculating:

$$\text{Torque} = \text{Force} \times \text{The distance of the force from center of rotation}$$

$$= 10 \text{ lbs.} \times 2 \text{ feet}$$

$$= 20 \text{ lb. - feet}$$

Both dc motors and dc generators rely upon commutator action for their operation. Remember that a commutator reverses the direction of current flow in an armature for each 180° of rotation. Therefore, for the same reasons that a commutator changes ac to pulsating dc in a generator, it also converts dc to ac in a motor.

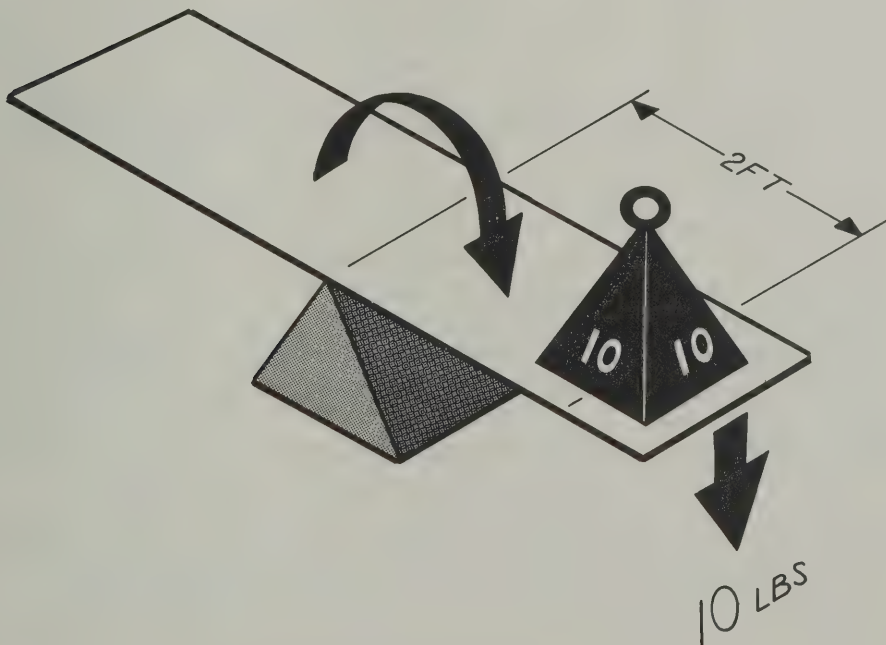


Figure 21

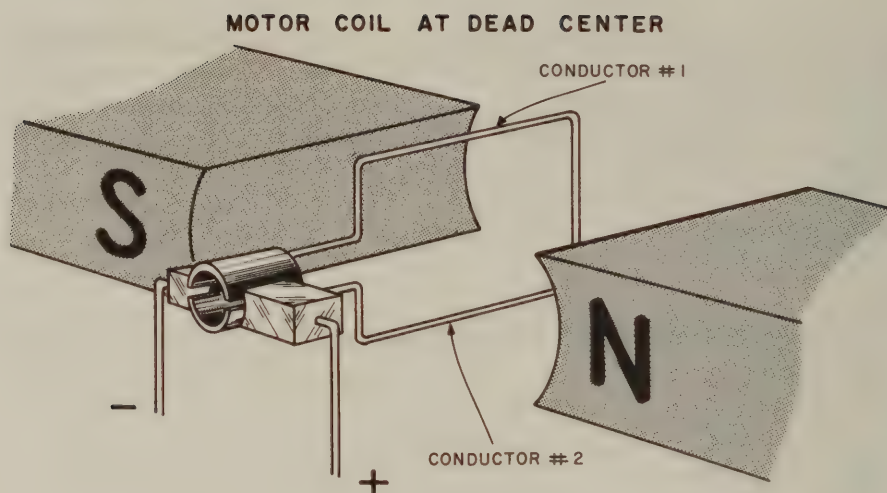


Figure 22

By adding a commutator to the simple, one-loop armature, also called the **ROTOR** in a motor, we find that continuous rotation occurs. A direct current is applied to the brushes with the polarity shown in Figure 20. The current is applied to the rotor through the commutator segments. The resulting interaction of fields causes conductor No. 1 to be forced upwards, and conductor No. 2 to be forced downwards as explained by the motor principle. Rotation in this direction continues until the rotor reaches **DEAD CENTER**, as shown in Figure 22. Dead center occurs when the brushes touch both commutator segments simultaneously. When the rotor reaches dead center, note that the coil and the dc supply are momentarily shorted out. The rotor continues to rotate through dead center, however, due to its inertia. That is, since it has already been rotating, the rotor has enough momentum to keep rotating for a short time even though the dc supply is momentarily shorted out. As the coil rotates past dead center, the commutator segments reverse the direction of current through the coil. Thus, conductor No. 1, which was being forced up through the stationary field, is now being forced down through the field. Similarly, conductor No. 2 which was being forced down is now being forced up. The overall result is that the rotor continues to turn in the same direction. This process is repeated each 180° of rotation. Torque is exerted on the rotor at all times except at dead center.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

DC GENERATORS

SELF-EXCITED GENERATOR

REGULATION

BASIC CONSTRUCTION

ARMATURE

MOTORS

THE MOTOR PRINCIPLE

TORQUE

9. The electromagnets which provide the flux in a generator are called _____.
10. The excitation circuit in a self-excited generator is completely independent of the armature and the load circuit. True or False?
11. Series generators are used in _____ applications.
12. In a four-pole generator, the N and S pole pairs are placed opposite each other. True or False?
13. The distance between two adjacent stator poles is called a (a) fractional pitch, (b) pole pitch.
14. The right hand rule for motors is used to determine the direction of (a) current, (b) conductor motion.
15. If the N and S poles of the stationary magnetic field shown in Figure 19 were the reverse of the positions shown, the (a) left hand generator rule would be used, (b) the conductor would be forced down.
16. In a motor, the commutator (a) produces ac in the armature coil, (b) converts the armature ac to pulsating dc.
17. When a dc motor armature reaches dead center, the commutator reverses current direction through (a) the load, (b) the armature coil.

DC MOTORS

A dc motor, like a dc generator, consists of end bells, frame, brush assembly, commutator assembly, armature assembly and pole pieces. The stationary field is not always a permanent magnetic field. It may be supplied by an electromagnetic field, as with a dc generator. Figure 23 shows the field windings and pole pieces of a dc motor. The field windings are excited from a dc power supply, and provide a steady electromagnetic field. The polarity of the field is determined by the direction of the current flow through the windings. Thus the polarity is adjusted such that one of the pole pieces becomes an N pole and the other pole piece (in a two-pole motor) becomes an S pole. This establishes the stationary field. However, a motor also requires dc excitation to the armature. This could be supplied by a separate dc power supply, but often the field windings (stator) and the armature assembly (rotor) are supplied from the same source. By adding more loops of wire to the armature, the motor principle is aided by REPULSION and ATTRACTION. Actually, the repulsion and attraction effect becomes the dominant factor in dc motor operation when many loops of wire are used. Therefore, we may consider this the principal method of operation.

Figure 23A shows a simplified illustration of a two-pole dc motor. With multiple loops of wire wound on an iron core, the current through the coil produces an electromagnet. Thus, the armature has an N and S pole. The N pole of an electromagnet can be determined by the left-hand current rule. To find polarity by using the left-hand current rule, the fingers of the left hand are pointed parallel to the windings in the direction of current flow through the windings. The extended thumb then points to the N pole of the electromagnet. Remember that the strength of the electromagnetic field depends upon the number of turns in the windings and the current through the winding. Thus, the single-turn armature used in previous examples did not develop an appreciable magnetic field, and did not interfere with the analysis of forces in the demonstration of torque.

The armature polarity is shown in Figure 23A. Under these conditions the armature will rotate in the direction shown until the N pole of the armature is next to the S pole of the field windings. In other words, since unlike poles attract, the poles of the coil are attracted to the unlike poles of the field windings. Figure 23B shows what happens when the coil reaches the commutating (neutral) plane. Note that just as the coil reaches the neutral plane and is carried past dead center, the commutator reverses the direction of current through the coil. This also reverses the polarity of the electromagnetic field about the coil. The end of the coil which was an N pole now becomes an S pole. Since the coil has already been carried past dead center, the armature is REPELLED by the field poles in the same direction that it was previously rotating. This process continues for each 180° of rotation, and the armature continues to turn in the same direction. The

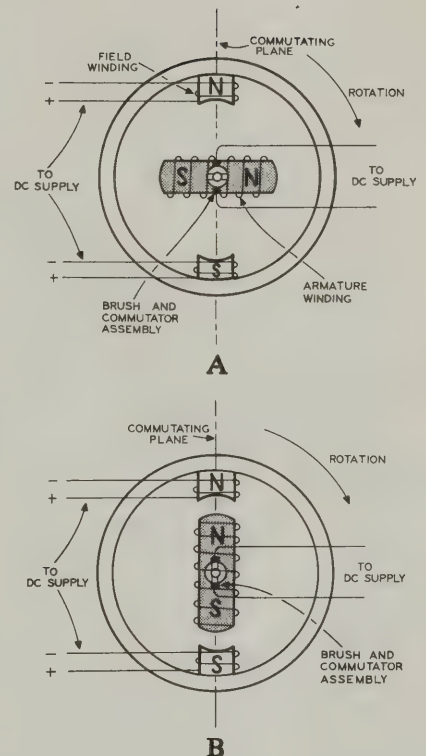


Figure 23

armature can be made to turn in the opposite direction by reversing the leads on either the field windings or the armature. If, however, the leads are reversed on BOTH the field and armature windings, the armature will turn in its original direction.

There are three types of dc motors—SERIES, SHUNT, and COMPOUND. This refers to the method in which the field coils and the armature coils are connected. It is important to know the characteristics of each type of motor in order to use them properly.

Series Motors

A series dc motor, as the name implies, is one in which the field coil and the armature coil are in series. As you can see in Figure 24, a basic series motor diagram is the same as a series generator diagram. Since the starting current through the armature also flows through the field coils, the series motor develops a high starting torque, but has poor speed regulation. Figure 24B shows a typical series motor characteristic curve. Note that as the speed decreases, the amount of torque delivered to the load increases. The reason for this is that as the armature slows down, the CEMF developed in the armature decreases, the current through the armature and field

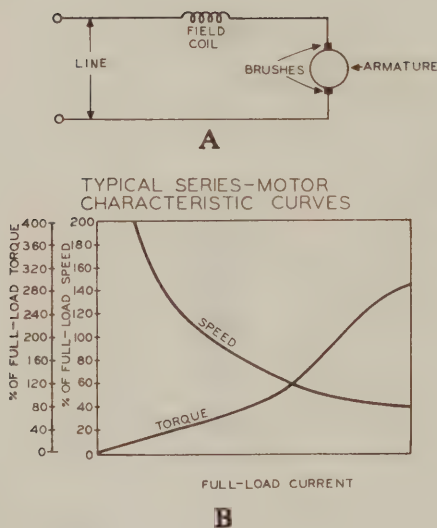
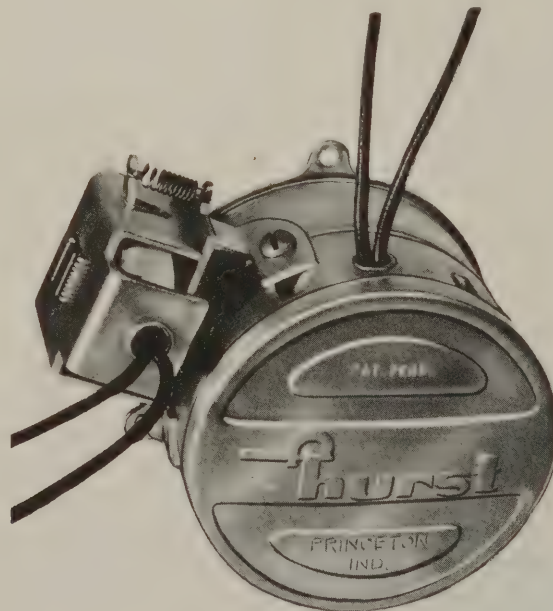


Figure 24



This motor uses a separately controlled actuating mechanism and a clutch-brake assembly to permit fast start-stop operation.

Courtesy Hurst Manufacturing Corp.

increases and the torque increases. Remember that increasing the current through the field coils and the armature increases their fields which causes the increased torque. This means that a series motor runs fast with a light load, and runs slower as the load is increased. In fact, if a series motor is allowed to run without a load, it may run so fast that the armature will fly apart. It is for this reason that series motors should not be used with belt coupling to the load. If the belt were to break, the motor would be without a load, and could destroy itself.

Shunt Motors

Shunt dc motors, again similar to shunt generators, are connected with the field windings and the armature windings in parallel. In this case, the field windings carry only part of the excitation current, and are made of smaller-gauge wire than those of a series motor. As shown by the characteristic curves in Figure 25, a shunt motor has good speed regulation, but low starting torque. This is just the reverse of the series motor characteristics. A shunt motor should not be used to start heavy loads. It may, however, be allowed to run with no load because speed depends primarily upon the resistance of the field windings. By adding a variable resistance to the field coil circuit (Figure 25C), motor speed can be controlled above or below its nominal value. Increasing the field resistance decreases the field current and the CEMF induced in the armature. This causes greater current to flow through the armature, increasing the torque and the motor will speed up. Decreasing the field resistance causes the field current to increase, increasing the CEMF developed in the armature, thus reducing armature current. The reduced current produces less torque and the motor slows down. Thus, the speed of a shunt dc motor is relatively independent of the load.

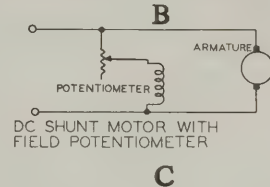
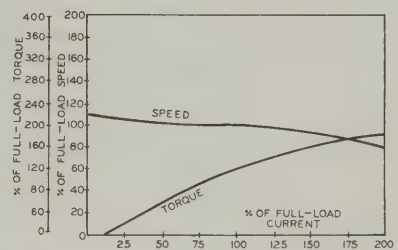
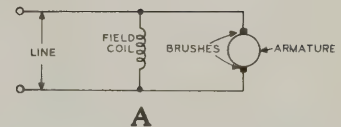


Figure 25

Compound Motors

A compound motor uses both series and shunt field windings. The advantage of this arrangement is that the compound motor can be made to produce a variety of operating characteristics. By adjusting the polarity and placement of the field windings with respect to each other, a number of motor types can be designed. For example, when the series and shunt fields are in series-aiding, the motor is a CUMULATIVE COMPOUND motor. When the fields are series-opposing, the motor is a DIFFERENTIAL COMPOUND motor.

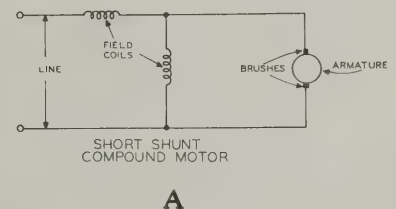
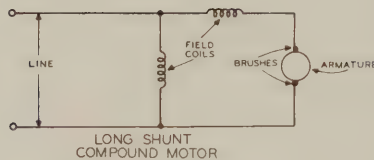


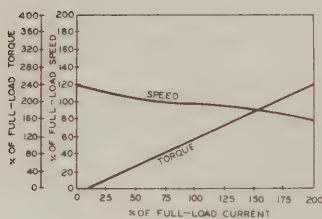
Figure 26

Another factor in the design of a compound motor is the placement of the series winding. Figure 26 shows this relationship. When the shunt winding is connected directly across the armature, the motor is a SHORT SHUNT compound motor.



B

When the shunt winding is connected in parallel with the ARMATURE AND THE SERIES WINDING, the motor is a **LONG SHUNT** compound motor. This is shown in Figure 26B. By controlling the field strengths around the coils using cumulative/differential, or short/long shunt methods, a compound motor can produce any characteristic that can be produced by a pure series or shunt motor.



C

Figure 26

Figure 26C shows the characteristic curves which could be obtained for one type of compound motor. Remember that this is merely an example, and many types can be obtained. Note that the compound motor has a speed characteristic of the shunt motor, and also has a starting torque characteristic which is nearly as good as the series motor. In addition, a compound motor can be operated safely without a load.

Motor Starting

A motor armature offers very little resistance to a dc source when it is not rotating. In fact, a motor which normally draws only 10 amps while operating may draw 100 amps or more when the armature is at rest. The difference is due to the COUNTER EMF (CEMF) generated by the armature when it is rotating.

As we learned in the section on generators, a coil rotating in a magnetic field always produces a voltage by induction. A motor armature is no exception. It, too, produces an emf (when rotating) which OPPOSES the applied excitation voltage. However, when the motor is first turned on and is not yet rotating, there is no CEMF to oppose the high starting current. This starting, or INRUSH CURRENT, would easily burn out the armature windings if not limited to a safe value. Once the armature is rotating at or near normal speed, the CEMF becomes great enough to safely limit the source current.

A variety of motor starting techniques are used, ranging from manual, semi-automatic, to automatic methods. Figure 27A shows a simple manual starter. A potentiometer having a number of set points is placed in series with the armature coil. Before current is first applied to the armature, the potentiometer wiper arm is set to position No. 1 as shown. This places maximum resistance in the armature circuit and limits starting current. The armature turns slowly at this point and begins to build up a CEMF. When the armature has reached the maximum speed for this setting, the wiper arm may be moved to set point No. 2. This reduces the resistance in the armature circuit, and the armature picks up speed. This, in turn, produces a greater CEMF. Further reductions in resistance now will produce smaller current increases, and the wiper arm may be moved rather rapidly to the last

(minimum resistance) set point. Note also that as the resistance is removed from the armature circuit, it is added to the field coil circuit. This merely establishes the maximum speed of the motor.

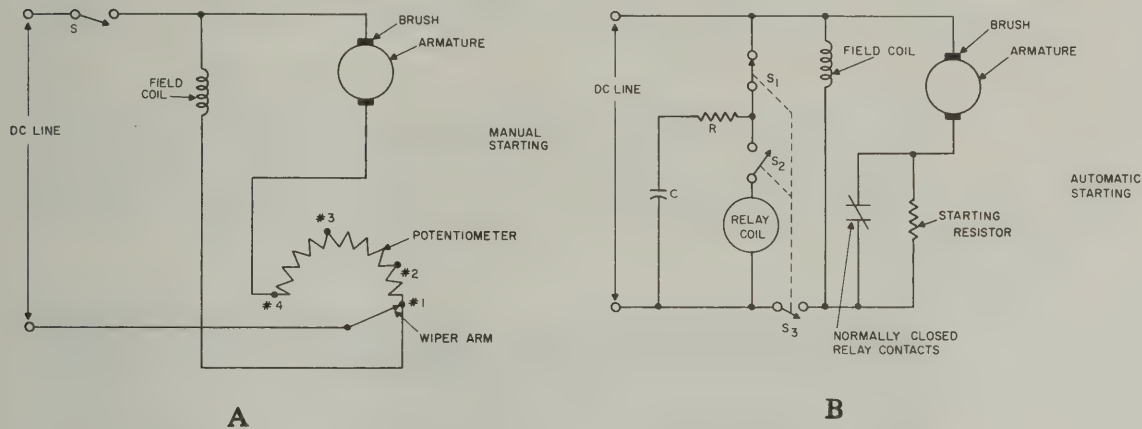


Figure 27

Another method uses a resistance-capacitance (RC) time delay circuit as shown in Figure 27B. Switches S_1 , S_2 , and S_3 are in the positions shown when the motor is off. However, capacitor C charges to the dc line voltage through R and S_1 . The switches are GANGED so that when the motor is started, S_1 opens while S_2 and S_3 close. When S_2 closes, the capacitor discharges through resistor R and the relay coil, thus energizing the relay. This opens the normally closed (NC) relay contacts and places the STARTING RESISTOR in series with the armature winding. By properly choosing the values of R and C , the relay will remain energized long enough to permit the armature to build up speed. Since the starting resistor remains in the circuit during this time, current through the armature is limited to a safe value. When the capacitor has discharged, the relay deenergizes, and the contacts again close, shorting out the starting resistor. However, current through the armature remains at a safe level because the armature has gained sufficient speed to produce the required CEMF.

Motor Efficiency

The efficiency of a machine is defined as the ratio of its output power to its input power, expressed as a percent. This is stated in the formula:

$$\text{Efficiency} = \frac{P_o}{P_i} \times 100$$

In order to calculate efficiency, the input and output power must be in the same units. For example, if a system produces 75 watts of output power while requiring 100 watts of input power, the efficiency of the system is

$$E_{tt} = \frac{P_o}{P_i} \times 100 = \frac{75}{100} \times 100 = .75 \times 100, \text{ or } 75\%.$$

While a motor input power is measured in watts, the motor output is measured in horsepower (hp). Therefore, in order to calculate motor efficiency, it is necessary to convert horsepower to watts. Since 1 hp is equal to 746 watts, the output hp is multiplied by 746 to obtain the motor output in watts.

Motors are always less than 100% efficient due to power losses in the windings and the friction of moving parts within the motor. Evidence of this is seen in the fact that motors generate considerable heat, and fans are often attached to the shaft for cooling. Fan blades also reduce the efficiency of a motor due to air friction.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

DC MOTORS

SERIES MOTORS

SHUNT MOTORS

COMPOUND MOTORS

MOTOR STARTING

MOTOR EFFICIENCY

18. Repulsion and attraction occurs in a dc motor armature because current through the rotor is a constant dc. True or False?
19. As the speed of a series motor increases, the amount of torque delivered to the load _____.
20. A dc motor which should not be run without a load is the (a) shunt type, (b) series type, (c) compound type.
21. When the series and shunt fields of a compound motor are connected in series aiding, the motor is a (a) differential compound, (b) cumulative compound.
22. A motor armature presents the same opposition to a dc source when rotating as it does when at rest. True or False?

THREE-PHASE POWER

In the section on generators, Figure 5 shows a single-loop armature winding and the current waveform produced. This waveform is single-phase since only one armature coil is used. When more than one armature coil is used in an ac generator, the output is MULTIPHASED. Figure 28 shows a simple

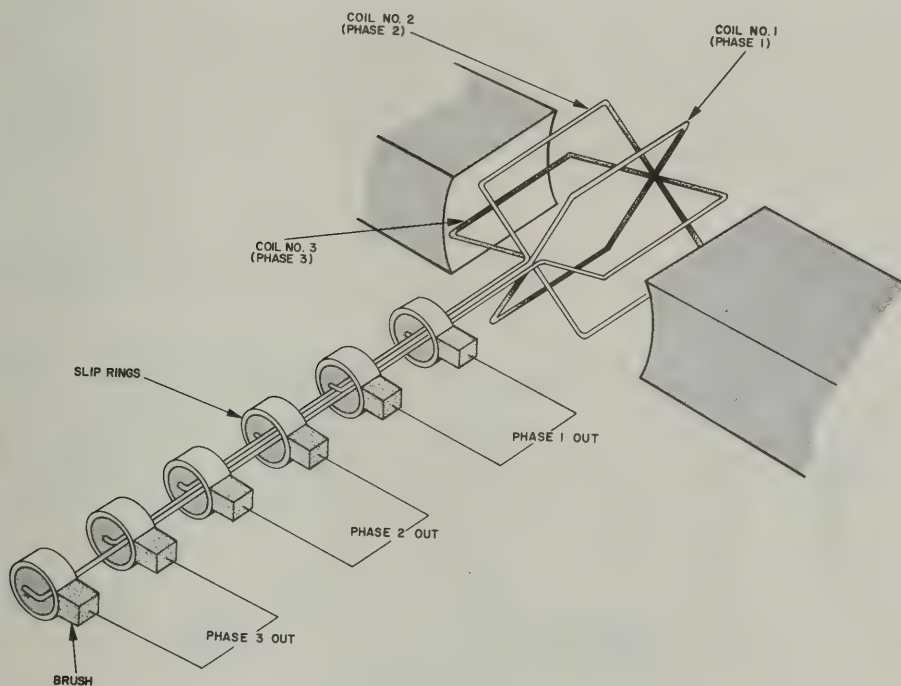


Figure 28A

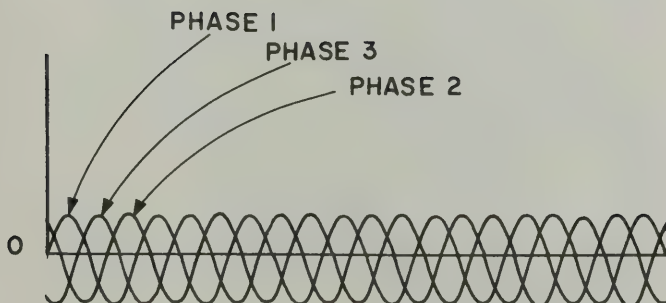


Figure 28B

three-phase alternator and the output current waveform. Note that the three sinewaves cross the baseline at different times, and also reach the peak values at successively later times. Three-phase power is widely used in industry for operating motors and machinery.

Synchronous AC Motors

Three-phase power can be used to operate a motor by producing a ROTATING MAGNETIC FIELD. This principle is shown in Figure 29A. The stator windings 1, 3, and 2 are connected to the incoming power phases 1, 3, and 2 respectively. Since each phase reaches peak amplitude at successively later times, the electromagnetic field becomes the strongest in each winding in succession. This creates the effect of an electromagnetic field with its N

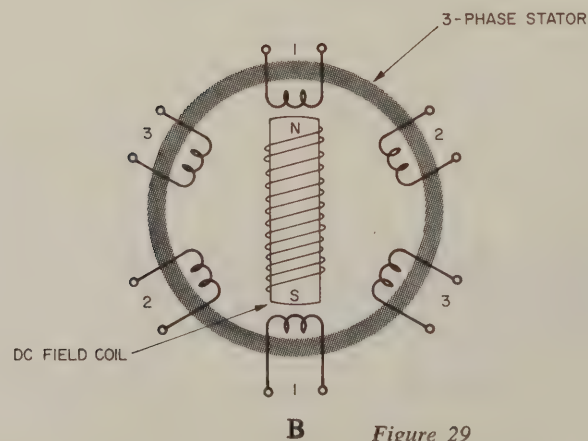
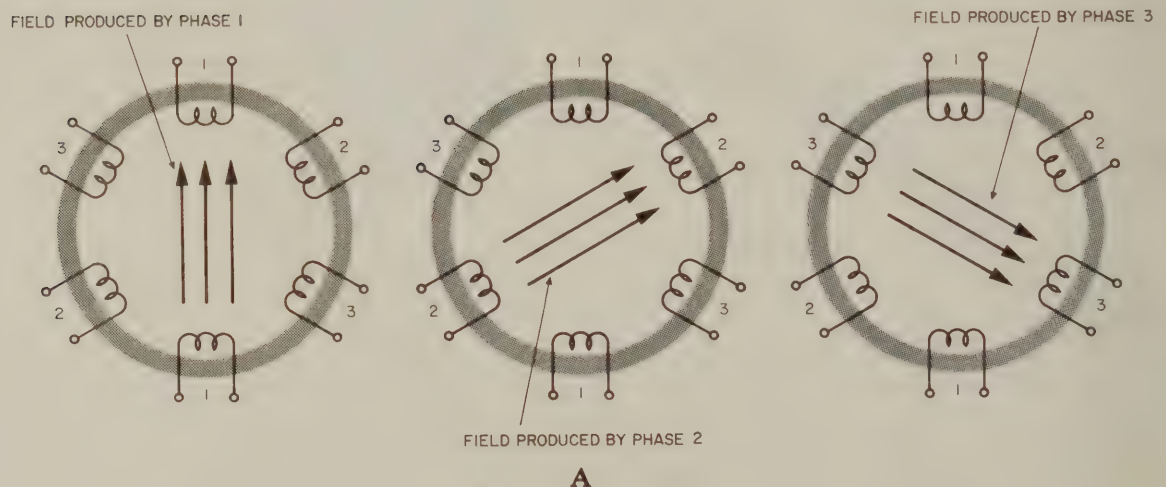


Figure 29

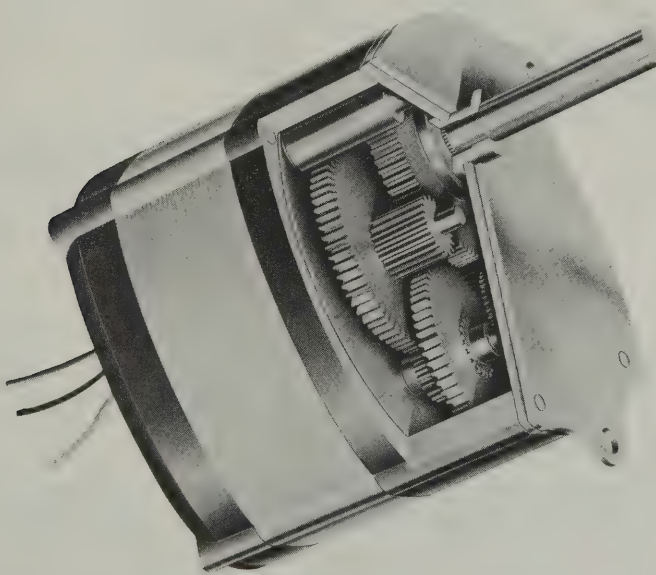
and S poles continually shifting around the stator hull. The field is, in addition, rotating one complete revolution for each cycle of phase 1, 2, or 3. This means that the field is rotating in SYNCHRONISM with the incoming line frequency. The rotor, shown as an electromagnet in Figure 29B rotates in synchroism with the stator field by maintaining alignment of its N and S poles with the S and N poles of the stator. Thus the rotor and the output shaft are also synchronized to the line frequency, and the motor rotates at SYNCHRONOUS SPEED. This is the only speed at which a synchronous motor can be used.

Note the similarity between the dc excited synchronous motor and the synchronous alternator, described previously, and illustrated in Figure 7. A dc excitation voltage may be applied through slip rings, or ac excitation may be applied through a commutator, in order to produce a STEADY dc field.

SYNCHRONOUS MOTORS, also called sync motors, are used in applications requiring precise speed control. Sync motors are used in electric clocks since they run in synchronism with the 60 Hz line frequency.

Three-Phase Induction Motors

A three-phase induction motor requires a revolving stator field derived from the three-phase power source just as does a sync motor. The main difference



Gear assemblies are used with sync motors to provide the exact output rpm required for a given application.

Courtesy Hurst Manufacturing Corp.

is that an induction motor does not require dc excitation in the rotor. Therefore, a basic induction motor uses no slip ring or commutator assembly. Instead, current is INDUCED in the rotor by the cutting of electromagnetic flux lines (from the rotating stator field) across the rotor windings. This voltage induced in the rotor causes a current to flow which sets up an electromagnetic field in the rotor. Thus, while the sync motor uses external power to supply the rotor current, the induction motor generates its own rotor current. This can be seen by examining one type of induction motor which uses a **SQUIRREL-CAGE** rotor.

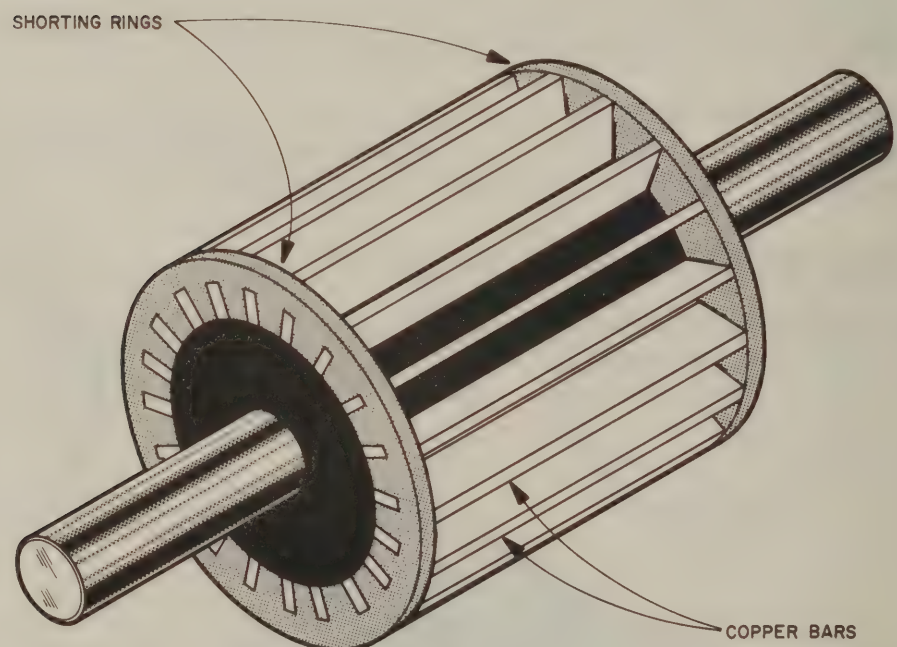


Figure 30

A squirrel-cage rotor is shown in Figure 30. This rotor uses heavy copper bars inserted at intervals in a drum which is attached to the rotor shaft. Shorting rings are used, as shown, to short the bars together. The rotor does not use wire windings as found in other types of motors. Instead, the copper bars serve as the windings, and are capable of carrying very high currents. The bars behave as though the rotor had a large number of turns in each coil while still having very low resistance. In operation, the rotating stator field cuts across the copper bars and induces voltage in them. Since they are shorted together, current flows through the bars and the shorting rings. This creates a polar electromagnetic field around each of the "coils" which react with the rotating field. Thus, the rotor turns.

Induction motors can never operate at synchronous speed. If the rotor were to turn at the same speed as the rotating field, no lines of force would be cut by the rotor conductors, and no rotor field would be produced. An induction motor must, therefore, rotate at something less than synchronous speed. The difference between the synchronous speed and the actual rotor rpm is called SLIP.

Slip is found by subtracting the rotor speed from the synchronous speed. The percentage of slip can be found by using the formula:

$$\% \text{Slip} = \frac{S_s - S_r}{S_s} \times 100$$

where S_s = the synchronous speed in rpm

and S_r = the actual rotor speed in rpm

For example, find the percentage of slip in an induction motor when the synchronous speed is 3,600 rpm and the motor speed at full load is 3,425 rpm.

$$\text{Slip} = S_s - S_r = 3,600 \text{ rpm} - 3,425 \text{ rpm} = 175 \text{ rpm}$$

$$\% \text{Slip} = \frac{175 \text{ rpm}}{3,600 \text{ rpm}} \times 100 = 4.8\%$$

The full load slip in most induction motors varies between 4 and 6 percent. Should an induction motor become heavily overloaded or stalled (the slip would be 100%), the rotor would be damaged. In general, slip in an induction motor should not exceed 10%.

Induction motors also have wirewound rotors. In this case, the coil ends are shorted together, and the operation of the motor is the same as for the squirrel cage rotor.

SINGLE-PHASE INDUCTION MOTORS

Most home and business appliances operate on single-phase ac power. For this reason, single-phase ac motors are in widespread use. Furthermore, induction motors account for a large percentage of these, due to their rugged construction, maintenance-free operation, and low cost.

A single-phase induction motor operates on the principle of induction, just as does a three-phase motor. However, in a single-phase motor, the stator field does not rotate, but instead merely alternates poles as the single sinewave voltage swings from positive to negative. Because of this, the stator field remains lined up in one direction with the poles changing position once each cycle. As a result, a single-phase induction motor will run once it has been started, but it has no means for starting by itself. By comparing Figures 29 and 31, one can see that a true spinning motion is not obtained from a single-phase motor as it is in a three-phase motor.

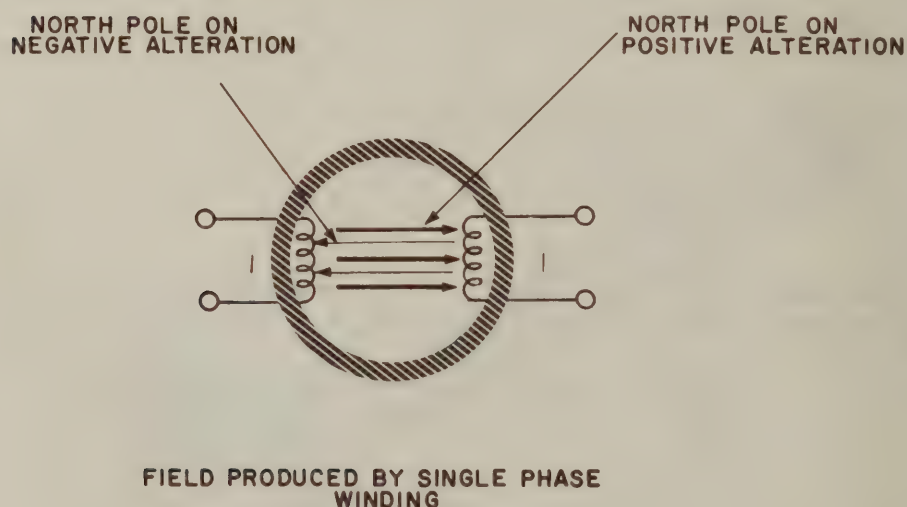


Figure 31

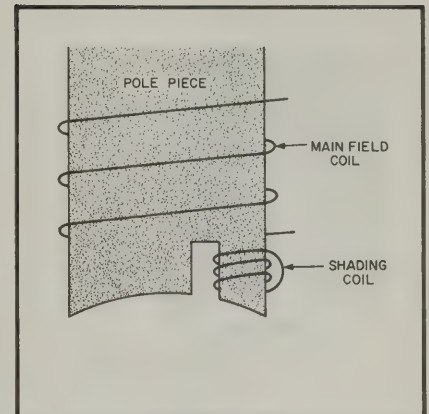
When a single-phase motor is first turned on, the rotor may not be aligned perfectly with the unidirectional stator field. If not, then a small voltage will be induced in the rotor, and it will tend to move into alignment. (This is similar to the dead-center position of a dc motor. Therefore, we will call the alignment position "dead center".) When the rotor reaches dead center, voltage is no longer induced because no lines of force are cut in the stator field. However, when the rotor is already spinning with some speed, inertia carries it past dead center, just as with a dc machine, and voltage continues to be induced in the rotor. In this way, the rotor turns by being attracted to the continually alternating stator poles.

A single-phase motor could be started by mechanically spinning the rotor, and then quickly applying power. However, most of these motors use some type of automatic starting.

Single-phase induction motors are often referred to by the starting method used. We will examine three of the most common types.

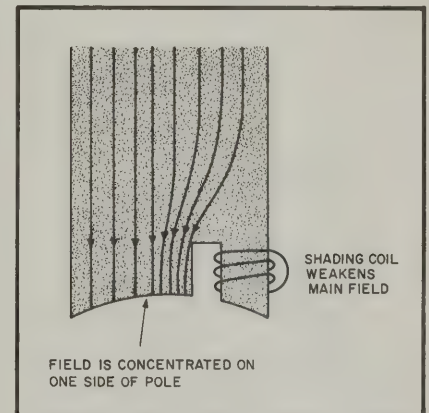
Shaded-Pole Motors

A shaded-pole motor is so named because it uses a small shorted **SHADING COIL** wound in a small notch in the stator pole piece. This is shown in Figure 32. In some motors this coil consists of a single copper ring or copper band. When the electromagnetic field builds up around the main coil, the flux cuts across the conductors of the shading coil. Since the shading coil is shorted, current flows which produces a field **OPPOSITE** that of main field. The main field is then strongest on the side away from the shading coil, as shown in Figure 32B. However, the field through the shading coil reaches maximum intensity much later—at a time when the main field is already decreasing. The electromagnetic field in the pole piece then appears to be stronger on the side nearest (or through) the shading coil as shown in Figure 32C. As you can see, this produces a sweeping motion from side to side in the stator pole piece. The direction of this field motion is shown in Figure 32D. Although small, this motion is enough to maintain an induced voltage in the rotor and start the rotor turning.



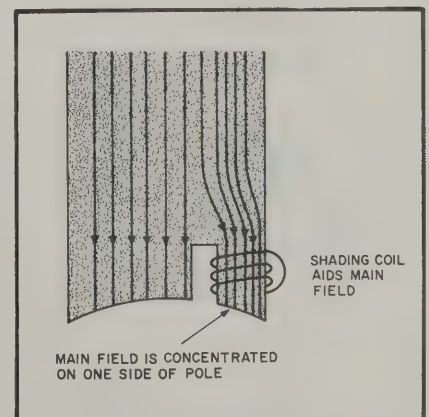
SHADED-POLE

A



ACTION OF SHADING COIL AS FIELD STRENGTH INCREASES

B



ACTION OF SHADING COIL AS FIELD STRENGTH DECREASES

C

Split-Phase Inductor Motors

A split-phase inductor motor produces a rotating field in the same way that a two or three-phase motor does. A separate stator winding, called a **STARTING WINDING**, is included in the stator field. This winding has only a few turns of very light gauge wire. When the ac line is applied to both windings simultaneously, a field builds up much more quickly around the starting coil than around the main coil. That is, the field around the main windings **LAGS** the field around the starting coil. This produces a partially-rotating stator field, which starts the motor as shown in Figure 33.

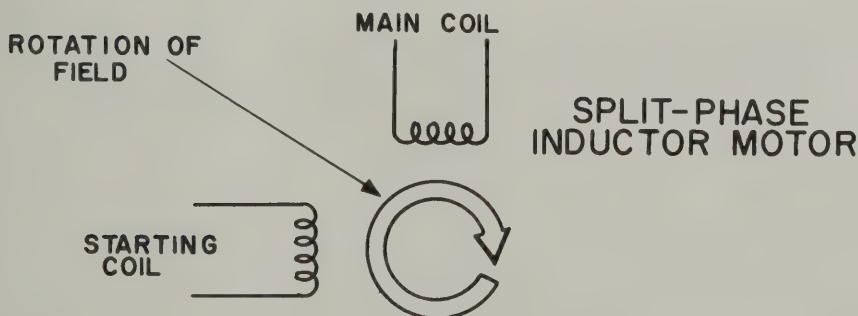
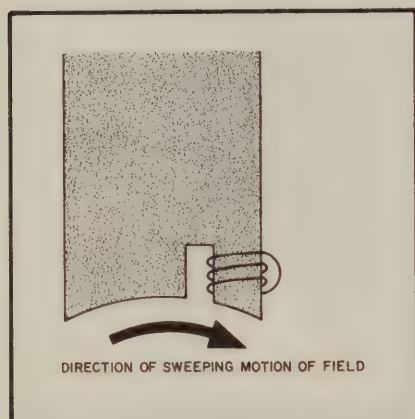


Figure 33

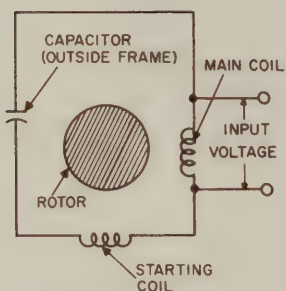
Figure 32



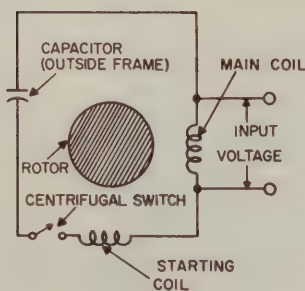
APPARENT MOTION OF FIELD
PRODUCED BY SHADING COIL

D

Figure
32



A



B

Figure
34

Once the motor is running, the starting winding must be removed from the circuit. Since the coil is light, continuous current through it would cause the winding to burn out. One widely-used method for doing this is to use a CENTRIFUGAL SWITCH physically mounted on the rotor, and wired in series with the starting coil. When the rotor gains sufficient speed to continue turning, using only the main field, centrifugal force operates the switch contacts, and breaks the starter coil circuit.

Split-Phase Capacitor Motors

Split-phase capacitor motors also produce a partially rotating stator field. In this case, however, the starting winding need not be a light coil. As shown in Figure 34, a large capacitor is placed in series with the second winding. When the line current is applied, a field builds up around the starting winding EARLIER than around the main winding. (The main winding field LAGS that of the starting winding.) Since the phase shift results from the capacitor, the starting coil may be identical to the main winding, and left in the circuit at all times. When this is done, the motor is called a **CAPACITOR-RUN motor**. If, instead, a lighter coil is used, and should not be left in the circuit, a centrifugal switch is used to break the series circuit after the motor has reached starting speed. This type of motor is called a **CAPACITOR-START motor**.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

THREE-PHASE POWER

SYNCHRONOUS AC MOTORS

THREE-PHASE INDUCTION MOTORS

SINGLE-PHASE INDUCTION MOTORS

SHADED-POLE MOTORS

SPLIT-PHASE INDUCTOR MOTORS

SPLIT-PHASE CAPACITOR MOTORS

23. A rotating field in a three-phase motor revolves at the same frequency as the incoming line voltage. True or False?
24. A three-phase induction motor requires current from an external source flowing in the rotor circuit. True or False?
25. The difference between the synchronous speed and the actual speed of an induction motor is called _____.
26. A single-phase induction motor requires a (a) rotating magnetic field, (b) separate starting technique.
27. A shading coil produces a (a) small sweeping field action, (b) rotating field.
28. Starting windings are used in (a) three-phase induction motors, (b) split-phase inductor motors.
29. A centrifugal switch is used to prevent damage to a (a) starting winding, (b) shading coil.
30. A motor which includes no means for removing the winding with the series capacitor is called a (a) capacitor start motor, (b) capacitor run motor, (c) inductor motor.

SUMMARY

Motors and generators are transducers which convert electric energy into mechanical energy, and mechanical energy into electric energy, respectively. Generators produce electricity by magnetoelectric induction. The voltage induced into any generator armature is ac. When this ac voltage is taken from the armature via slip rings and brushes, the generator is called an alternator. DC power can be obtained from a generator by using a commutator. Since the current in the armature reverses direction, a commutator maintains dc conduction by reversing the connections to the load every 180° .

All generators require dc excitation to the field windings, the dc excitation builds up a steady electromagnetic field which, when cut by the armature conductors, induces a voltage in the armature. When the excitation current is supplied from a separate dc source, the generator is separately excited. When the excitation current is supplied by the voltage induced into its own armature (rectified by the commutator), the generator is self-excited.

Output voltage regulation in generators is accomplished in several ways. In general, a sensing device placed in the load circuit varies the resistance in series with shunt field windings.

The major parts of a dc generator are the frame, armature assembly, commutator assembly, brush assembly, pole pieces, field coils, shaft, and end bells. The armature coils are wound in any of several ways, but all are designed to provide efficient operation. The coils are recessed in slots on a drum, and can be allowed to rotate very close to the stator poles without touching them. Common types of windings are lap winding and wave winding.

Motor construction is similar to generator construction. A basic dc generator can also be connected to operate as a dc motor. Motors operate on the principles of interaction between electromagnetic fields. This interaction produces torque which can be used to turn a motor shaft.

A motor armature consisting of a large number of wire turns acts as a strong electromagnet. Since the stator pole pieces are also electromagnets, a practical motor is made by causing the poles of the stator and the armature (rotor) to attract and repel each other.

The three types of dc motors are series, shunt, and compound. A series motor develops high starting torque, but has poor speed regulation. Also, a series motor will damage itself if allowed to run without a load. A shunt motor has good speed regulation, but low starting torque. It may, however, be run

without a load since speed depends almost entirely upon the resistance of the field windings. A compound motor combines the features of the series and shunt motors.

DC motors (other than very small dc motors) require a higher resistance in series with the armature during starting than when the motor is running. This is due to the counter emf produced by generator action when the motor is running. The CEMF limits the current through the armature when the motor is running, but CEMF is not present when the armature is at rest. Thus, the inrush current must be limited when the motor is first turned on. Manual and automatic methods are used for starting dc motors.

Three-phase ac motors use a revolving magnetic field in the field windings to turn the rotor. Three-phase power is convenient to use because the phases are separated by 120° , and can be applied separately to a different set of stator poles.

A synchronous motor is one in which the rotor turns at the same speed as the rotating field. Sync motors are used where precise speed control is required.

Induction motors do not require dc excitation current in the rotor. A simple induction motor, such as a squirrel-cage motor, has neither slip rings nor a commutator. Instead, the field about the rotor is produced by inducing a voltage in the rotor windings as the rotating field cuts the rotor conductors.

An induction motor must run at something less than synchronous speed in order to maintain an induced voltage in the rotor. The difference in speed between the rotating field and the rotor speed is called slip.

Induction motors will operate on single-phase power if some means is provided for starting them. Once started, single-phase motors will continue to run on the single-phase source. Thus, these motors are named for the particular starting method used.

A shaded-pole motor is a single-phase motor which creates a slight sweeping action in the stator poles which is enough to start the motor. Split-phase inductor and split-phase capacitor motors use separate starting windings to create a partially-rotating stator field.

IMPORTANT DEFINITIONS

ALTERNATOR — A magnetoelectric generator for producing alternating currents.

ARMATURE — In a generator, the windings into which a voltage is induced. In a motor, the windings which exert torque in the presence of a magnetic or electromagnetic field. In a dc machine, the armature is the rotor. In ac machines the armature may be either the rotor or the stator.

CAPACITOR-RUN MOTOR — A split-phase capacitor motor in which the winding in series with the capacitor is used for starting and then remains in the circuit during motor operation.

CAPACITOR-START MOTOR — A split-phase capacitor motor in which the starting winding is removed from the circuit during motor operation.

COMMUTATOR — An assembly consisting of brushes and circular SEGMENTS which effects a reversal of current direction in the rotating member of a motor or generator.

COMMUTATING PLANE — In a dc generator, the plane through a circular field in which lines of force are not cut by the rotating member.

DEAD CENTER — In a dc motor, the point at which the brushes short out two or more commutator segments, momentarily, thus cancelling the exertion of torque at that point.

DYNAMO — A machine capable of converting mechanical energy to electric energy by magnetoinduction.

ELECTROMAGNET — An electromechanical device which produces an electromagnetic field by the action of current flowing through a coil of wire.

FLASHING THE FIELD — In starting a self-excited generator, the field windings may require external excitation briefly during the first few revolutions in order to build up the electromagnetic field.

FIELD EXCITATION — The current supplied to the field windings of a motor or generator, thus producing an electromagnetic field.

FRACTIONAL PITCH WINDING — An armature winding in which the distance between the coil sides does not coincide with the distance between stator poles.

GENERATOR — A magnetoelectric mechanism which produces an electric current output as a result of rotational mechanical input force.

IMPORTANT DEFINITIONS (Continued)

LAP WINDING — An armature winding arrangement in which the ends of each coil are connected to adjacent commutator segments.

LEFT-HAND CURRENT RULE — A rule for determining the direction of the lines of force around a current-carrying conductor. When the fully-extended thumb of the left hand is pointed in the direction of current flow, the curled fingers of the left hand indicate the direction of the flux lines.

LEFT-HAND GENERATOR RULE — A rule for determining the direction of current generated by a conductor moving through a magnetic field. When the thumb, middle, and index fingers of the left hand are held at right angles to each other (the thumb in the direction of motion; the index finger in the direction of the field N to S), the middle finger indicates the generated current flow direction.

LONG SHUNT — A compound dc motor field coil arrangement in which the shunt winding is in parallel with a series circuit consisting of the armature and the series winding.

MAGNETOELECTRIC INDUCTION — The induction of an electromotive force in a conductor when relative motion is created between the conductor and a magnetic flux field.

NEUTRAL PLANE — See COMMUTATING PLANE.

POLE PITCH — The distance between two adjacent poles in the stator of a motor or generator (also called pole span).

RESIDUAL MAGNETISM — A small amount of magnetism which remains in an electromagnetic pole piece after the excitation current has been removed. This is due to the tendency of the electromagnet to permanently magnetize the core or pole piece.

RIGHT HAND RULE FOR MOTORS — A rule for determining the direction of motion resulting from the torque exerted on a current-carrying conductor in a magnetic field.

ROTOR — An assembly consisting of the rotating members of a motor or generator, including the input and/or output shaft.

SELF-EXCITATION — A method for supplying dc excitation to the field coils of a generator by feeding part of the armature output back to its own field windings.

SHORT SHUNT — A compound dc motor field coil arrangement in which the shunt winding is shunted directly across the armature.

IMPORTANT DEFINITIONS (Continued)

SLIP — The difference between the synchronous speed and the actual rotor speed of an induction motor.

SLIP RING — A metal ring attached to the rotating shaft of a motor or generator, thus permitting electrical connection to the rotor through BRUSHES.

SQUIRREL CAGE MOTOR — An induction motor so named because of the similarity between its rotor construction and a squirrel cage.

STARTING WINDING — An auxiliary field winding in a single-phase motor used with a phase-shifting component to create a partially-revolving field.

STATOR — The stationary portion of a motor or generator which includes the stator windings.

SYNCHRONOUS MOTOR — An ac motor in which the rotor speed turns in synchronism with the frequency of the applied field current.

TORQUE — A turning force, or moment of force, with respect to a particular point of origin.

TRANSDUCER — A device capable of converting one form of energy into another form of energy.

TURBINE — A device which transmits fluid-type mechanical energy into rotational mechanical energy.

WAVE WINDING — An armature winding arrangement in which the coil ends are connected to commutator segments two pole spans apart.

PRACTICE EXERCISE SOLUTIONS

1. transducer.
2. dynamo
3. True — The rotating armature of Figure 5 produces a negative and positive voltage peak during one revolution.
4. False — The MAXIMUM voltage is produced when the armature moves at right angles. Smaller voltages are induced at all other times except when the coil moves parallel to the field, at which time the induced voltage is zero.
5. 60 Hz. — Using the formula

$$f = \frac{p \times \text{rpm}}{120}; \quad f = \frac{2 \times 3600}{120} = 60 \text{ Hz}$$

6. True
7. commutating plane, or neutral plane.
8. False — The current in the armature is always ac. The commutator reverses the connections to the load.
9. field windings.
10. False — The excitation circuit includes the armature and the load.
11. constant current
12. False — The N and S pole pairs are placed next to each other.
13. (b) pole pitch.
14. (b) conductor motion.
15. (b) the conductor would be forced down.
16. (a) produces ac in the armature coil.
17. (b) the armature coil.
18. False — The current periodically reverses direction.
19. decreases.

PRACTICE EXERCISE SOLUTIONS (Continued)

- 20. (b) series type.
- 21. (b) cumulative compound.
- 22. False — A motor armature offers less opposition to a dc source when at rest. A rotating armature generates a CEMF which opposes the source.
- 23. True
- 24. False — The current flowing in the armature is the result of an induced voltage rather than from an external source.
- 25. slip.
- 26. (b) separate starting technique.
- 27. (a) small sweeping field action.
- 28. (b) split-phase inductor motors.
- 29. (a) starting winding.
- 30. (b) capacitor run motor.



ONE OF THE

BELL & HOWELL SCHOOLS

1106A

EXAMINATION
CHECK SHEET

1106A

1. A - TO GENERATE A VOLTAGE BY INDUCTION, THE ONLY REQUIREMENT IS THAT -- THERE BE RELATIVE MOTION BETWEEN A MAGNETIC FIELD AND AN ELECTRICAL CONDUCTOR. Magnetic lines of force can be cut by a conductor regardless of whether the field moves or the conductor moves.

2. C - A SELF-EXCITED SHUNT DC GENERATOR SHOULD BE OPERATED -- IN CONSTANT VOLTAGE APPLICATIONS.

The output voltage of a shunt generator is very sensitive to load changes. These generators are normally operated within relatively small ranges of output voltage.

3. D - A VOLTAGE REGULATOR IN A SHUNT DC GENERATOR RESTORES THE OUTPUT VOLTAGE -- BY VARYING THE CURRENT THROUGH THE FIELD WINDINGS.

The sensing element in a regulator is connected across the generator output. A variable resistance element of some type is placed in series with the field windings.

4. C - IN FIGURE 27A, THE STARTING RESISTANCE WHICH IS REMOVED FROM THE ARMATURE CIRCUIT IS ADDED TO THE FIELD COIL CIRCUIT. AFTER THE MOTOR HAS BEEN STARTED, THIS RESISTANCE ESTABLISHES THE MAXIMUM MOTOR SPEED BECAUSE -- THE MOTOR IS CONNECTED IN A SHUNT CONFIGURATION.

The resistance of the shunt field coil, and all resistances in series with the coil, establish the maximum speed of a shunt motor.

5. B - WHAT IS THE EFFICIENCY OF A 110 VOLT MOTOR WHICH DRAWS 20 AMPS AND DELIVERS 2.5 HORSEPOWER TO THE LOAD? -- 84.7%.

Using the formula:

$$E_{ff} = \frac{P_o}{P_i} \times 100$$

where

$$P_o = 2.5 \text{ hp} \times 746 \text{ watts/hp} = 1865 \text{ watts}$$

$$P_i = 110 \text{ volts} \times 20 \text{ amperes} = 2200 \text{ watts}$$

substitute

$$E_{ff} = \frac{1865 \text{ watts}}{2200 \text{ watts}} \times 100 = 84.7\%$$

6. D - IN THREE-PHASE MOTORS, THE THREE-PHASE POWER SOURCE IS CONNECTED TO -- THE STATOR WINDINGS.

A three-phase motor produces a revolving electromagnetic field around the stator hull. This is accomplished by phase differences which reach peak amplitude at successively later times.

7. D - INDUCTION MOTORS OBTAIN DC POWER IN THE ROTOR THROUGH -- INDUCED VOLTAGE. Relative motion is produced between the rotating or alternating field and the rotor windings. Thus, a voltage is induced in the rotor.

8. D - IF A TWO-POLE SYNCHRONOUS ALTERNATOR, TURNING AT 3600 RPM DRIVES AN INDUCTION MOTOR TURNING AT 3420 RPM, THE SLIP IS -- 5.0%.

Since the synchronous (input) frequency to the motor is 3600 rpm, the slip can be calculated from the formula

$$\% \text{ Slip} = \frac{S_s - S_r}{S_s} \times 100$$

$$\text{Slip} = 3600 - 3420 = 180$$

$$\% \text{ Slip} = \frac{180}{3600} \times 100 = 5\%$$

9. A - THE DIFFERENCE BETWEEN A THREE-PHASE INDUCTION MOTOR AND A CAPACITOR-START SINGLE-PHASE INDUCTION MOTOR IS THAT -- THE SINGLE-PHASE FIELD DOES NOT REVOLVE.

A single phase inductor motor produces an alternating, pulsating, sweeping, or partially-rotating field.

10. D - A SHADED-POLE MOTOR IS -- A SINGLE-PHASE INDUCTION MOTOR.

A shaded pole motor derives its name from the starting method used.

1106A

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1106A

Example: An earthworm has

A	☐
B	☐
C	☐
D	☒

A. two legs. B. four legs. C. six legs. D. no legs.

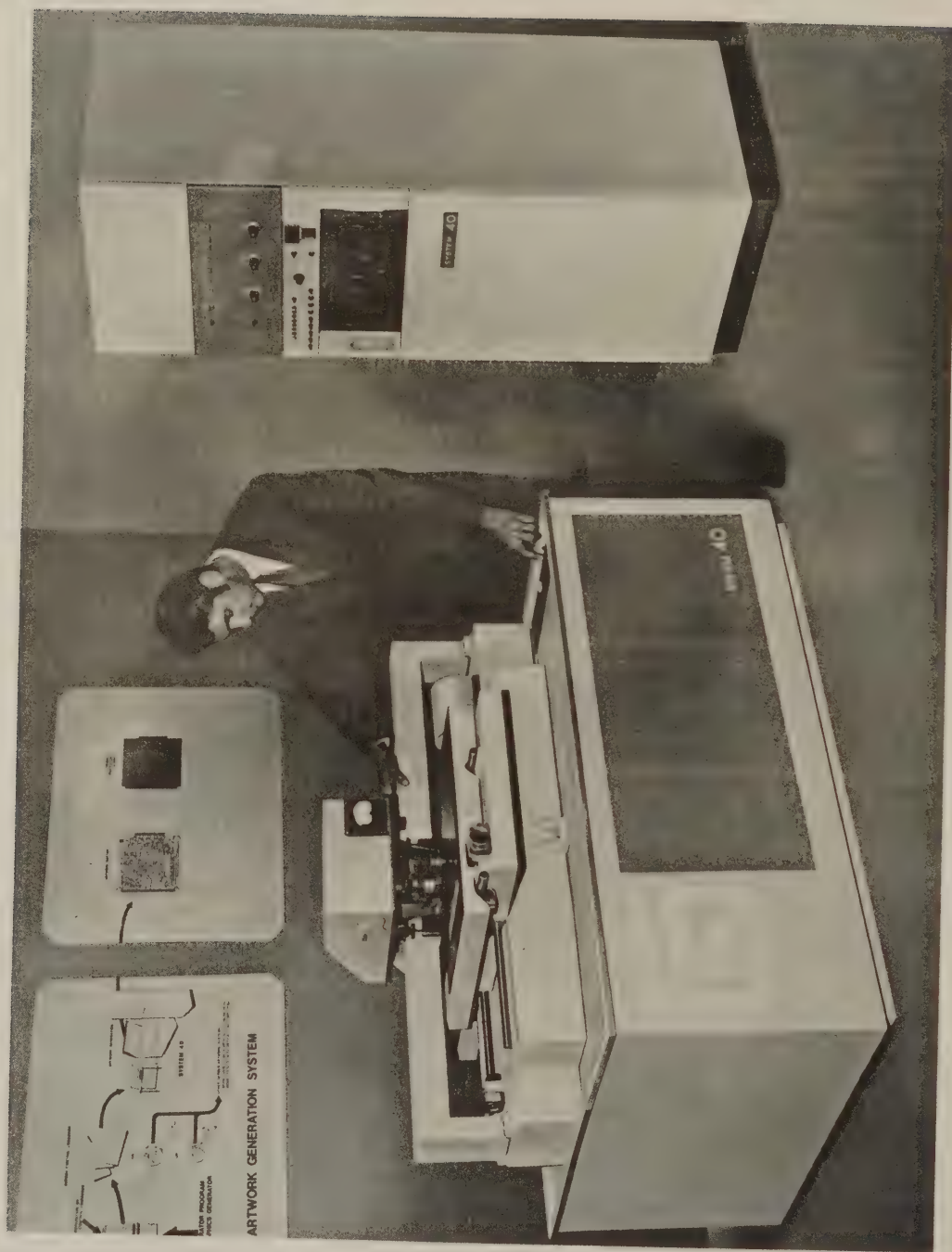
1. A ☒ To generate a voltage by induction, the only requirement is that
 B ☐ (A) there be relative motion between a magnetic field and an electrical conductor.
 C ☐ (B) a conductor move through a stationary magnetic field. (C) a magnetic field
 D ☐ move across a stationary conductor. (D) a conductor move at right angles to a
 magnetic field.
2. A ☐ A self-excited shunt dc generator should be operated
 B ☐ (A) in constant current applications. (B) below its rated voltage level. (C) in
 C ☒ constant voltage applications. (D) above its rated current level.
 D ☐
3. A ☐ A voltage regulator in a shunt dc generator restores the output voltage
 B ☐ (A) by varying the resistance of the load. (B) by varying the armature speed.
 C ☐ (C) by increasing the voltage across the armature. (D) by varying the current
 D ☒ through the field windings.
4. A ☐ In Figure 27A, the starting resistance which is removed from the armature circuit
 B ☐ is added to the field coil circuit. After the motor has been started, this resistance
 C ☐ establishes the maximum motor speed because
 D ☒ (A) the inrush current determines the speed. (B) the motor is connected in a series
 configuration. (C) the motor is connected in a shunt configuration. (D) the load
 determines the CEMF which limits the current through the armature.
5. A ☐ What is the efficiency of a 110 volt motor which draws 20 amperes and delivers
 B ☒ 2.5 horsepower to the load?
 C ☐ (A) 33.9%. (B) 84.7%. (C) 88.0%. (D) 40.2%.
 D ☐
6. A ☐ In three-phase motors, the three-phase power source is connected to
 B ☐ (A) the rotor windings. (B) the commutator. (C) slip rings. (D) the stator windings.
 C ☒
 D ☐
7. A ☐ Induction motors obtain ac power in the rotor through
 B ☐ (A) a separate dc power supply. (B) slip rings. (C) a commutator. (D) induced
 C ☐ voltage.
 D ☒
8. A ☐ If a two-pole synchronous alternator, turning at 3,600 rpm, drives an induction
 B ☐ motor turning at 3420 rpm, the slip is
 C ☐ (A) 9.8%. (B) 1.5%. (C) 98.0%. (D) 5.0%.
 D ☒
9. A ☒ The difference between a three-phase induction motor and a capacitor-start single-
 B ☐ phase induction motor is that
 C ☐ (A) the single-phase field does not revolve. (B) the single-phase motor uses a
 D ☐ commutator. (C) the three-phase motor cannot start by itself. (D) the single-phase
 motor does not rely upon an induced voltage for its operation.
10. A ☐ A shaded-pole motor is
 B ☐ (A) a synchronous motor. (B) a three-phase induction motor. (C) a compound
 C ☐ motor. (D) a single-phase induction motor.
 D ☒

PRINTED CIRCUIT TECHNIQUES

1109

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





Printed Circuit board manufacturers often use highly specialized equipment in order to produce the greatest number of quality boards at the lowest cost. This automatic artwork generator system (sometimes called a photoplottter) uses a computer to control a beam of light focused on a photosensitive material. These systems greatly reduce the time required to produce artwork masters for printed circuit board production.

Courtesy Gerber Scientific Instrument Co.

CONTENTS

Point-to-point Wiring	Page 5
Copper Foil Techniques	Page 5
The Silk Screen Method	Page 7
The Photo-etch Method	Page 8
Circuit Layout	Page 13
Conductor Routing	Page 15
Plating	Page 19
Circuit Tracing Printed Circuit Boards	Page 20
Repairing Printed Circuit Boards	Page 22
Appendix	Page 30
Copper Etchant Reactions	Page 30

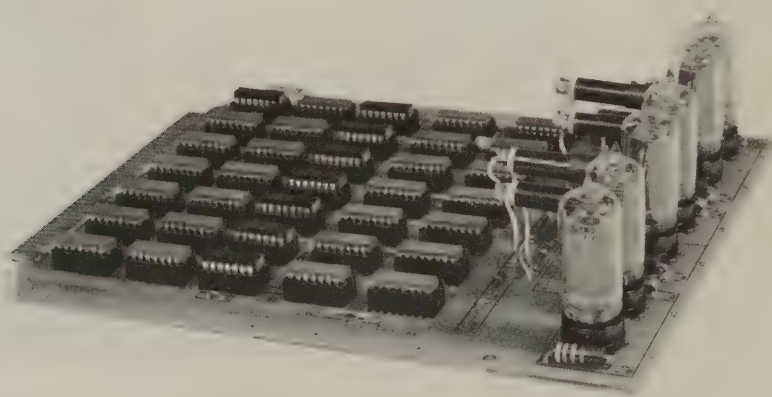
*Success doesn't happen. It is organized, pre-empted, captured
by concentrated common sense.*

—Frances E. Willard

PRINTED CIRCUIT TECHNIQUES

Printed circuits are largely responsible for the production of smaller and more reliable electronic equipment. In addition to the circuit operational characteristics, printed circuits are a part of what is known as **MINIATURIZATION**. Miniaturization has been developing steadily ever since the invention of the vacuum tube. Early vacuum tubes were large and consumed considerable power. Large operating voltages were also used. As tube technology advanced, tubes were made smaller, more efficient, and required lower operating potentials. This procedure led to the development of miniature and sub-miniature vacuum tubes.

Smaller and more sensitive components require special consideration in "packaging". That is, the methods used for connecting small devices together into a working unit are more critical than for large, nonsensitive elements. A system for obtaining a higher packing density (more components in a given area) was needed. To meet these requirements, printed circuits were developed which simultaneously solved two problems. One of these is the mounting configuration. Printed circuit boards provide a convenient base for mounting small components and allow many components to be mounted in a small area. The other advantage is in the interconnection function served by the printed circuits on the board. Thus, a printed circuit board holds the components and provides the interconnecting means for the components.



Circuit design flexibility and efficiency are enhanced using printed wiring boards. Here, integrated circuits, standard components, and readout-type vacuum tubes are mounted and interconnected on the same board.

Courtesy Chrono-Log Corporation

POINT-TO-POINT WIRING

Prior to printed circuits, all electronic equipment used **POINT-TO-POINT** wiring as the means for connecting the components together. This wiring is the hand-wired chassis type of circuitry. Components are soldered to sockets, tie-point terminal strips, etc. Conventional point-to-point wiring is still used where large components such as transformers and large electron tubes are included in the circuit. Printed circuits are often used in some parts of electronic equipment, while point-to-point wiring is used in others. Power supplies often use point-to-point wiring since the components such as transformers and filter capacitors are large and require considerable mounting space.

COPPER FOIL TECHNIQUES

Printed circuits are conductive-pattern boards which consist of a substrate and a very thin metal foil. The substrate is called the **LAMINATE**, and the foil is called the **CLADDING**. There are various types of laminates and cladding. Popular types of laminates are PHENOLIC PAPER, EPOXY PAPER, and EPOXY GLASS. Cladding is generally copper, although any conductive metal can be used.

There are several methods for obtaining the conductive pattern on a printed circuit board. Basic to all the methods is the preparation of the **ARTWORK**, which is a master pattern of the circuit. The artwork is either drawn directly on a copper-clad laminate, or produced on a copy board for later transfer to the copper surface. The simplest of the printed circuit fabrication techniques are the **PRINT-AND-ETCH** methods. In one of these methods, the artwork is drawn or taped up directly on the cladding. Figure 1 shows the use of a viscous **RESIST** material which is painted on the copper foil in the desired circuit pattern. When the complete pattern has been produced, the entire circuit board is immersed in a chemical bath. The chemical solution is called the **ETCHANT**, since all the copper which is not protected by the resist is etched away. The copper underneath the resist is not affected by the etchant since the resist is insoluble in the etchant. The result is a laminated board with copper remaining only where the resist material has been applied. When the resist is removed, using common solvents, the remaining copper is a conductive pattern which carries current in an actual circuit.

Instead of painting the resist on the cladding, special adhesive tape may be used to produce the artwork on the copper. This tape is also etch-resistant. As shown in Figure 2, the tape is applied directly to the copper using various shapes and lengths of tape. The tape is often pre-cut to standard widths for long conductor runs. Circular shapes are used which produce copper **PADS** on the circuit board. A small hole in the center of the circular pads indicates

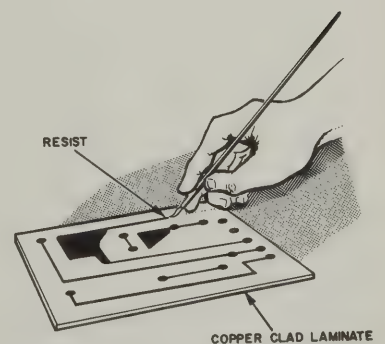


Figure
1

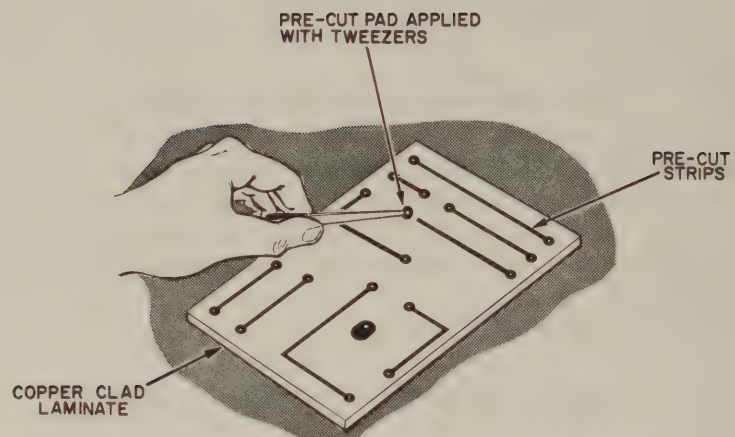


Figure 2

the center which is drilled through after the board has been etched. Thus, the pads and the center holes provide the means for mounting and connecting components to the board. The components are mounted on the side opposite the etched circuit, and the component leads are inserted through the board. This is shown in Figure 3. Since the copper paths are on the opposite side of the board as the components, they are shown by dashed lines in this figure. The enlarged pad area around the lead provides a base for a good solder connection to the lead.

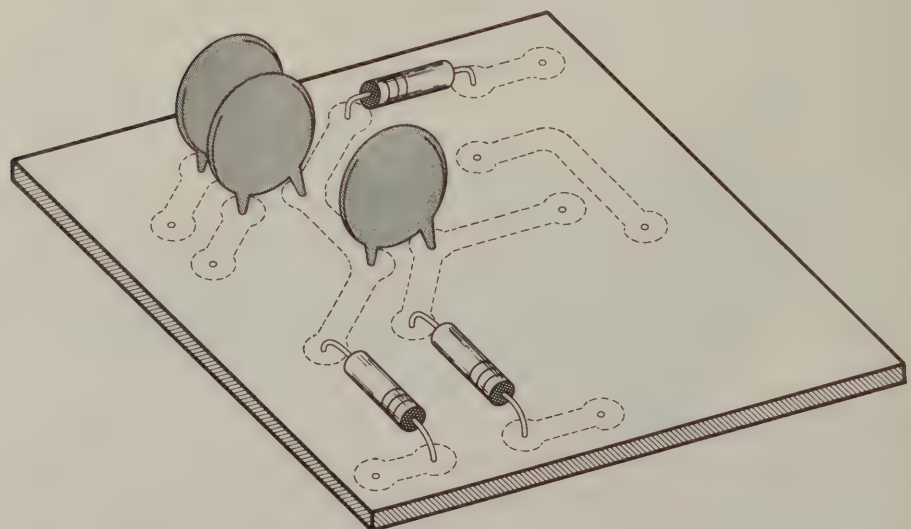


Figure 3

The Silk Screen Method

The two methods discussed so far are ~~simple manual techniques~~, and have certain limitations. Closely spaced, uniform, and neat-appearing circuit boards in production run volume are difficult and expensive to obtain using these methods. A faster print-and-etch method for applying the resist to the cladding is to use a silk screen, or stencil. Artwork in this technique is prepared by laying out the circuit pattern on a material suitable for use as a stencil overlay. The pattern in this technique is a reverse, or NEGATIVE, of the other methods. The overlay has cutout areas which correspond to the copper conductors. The resist material is screened onto the copper cladding, and the conductor pattern below the resist remains after the etch cycle. The screened technique is more suitable for producing large quantities of the same pattern since the overlay can be used repeatedly.

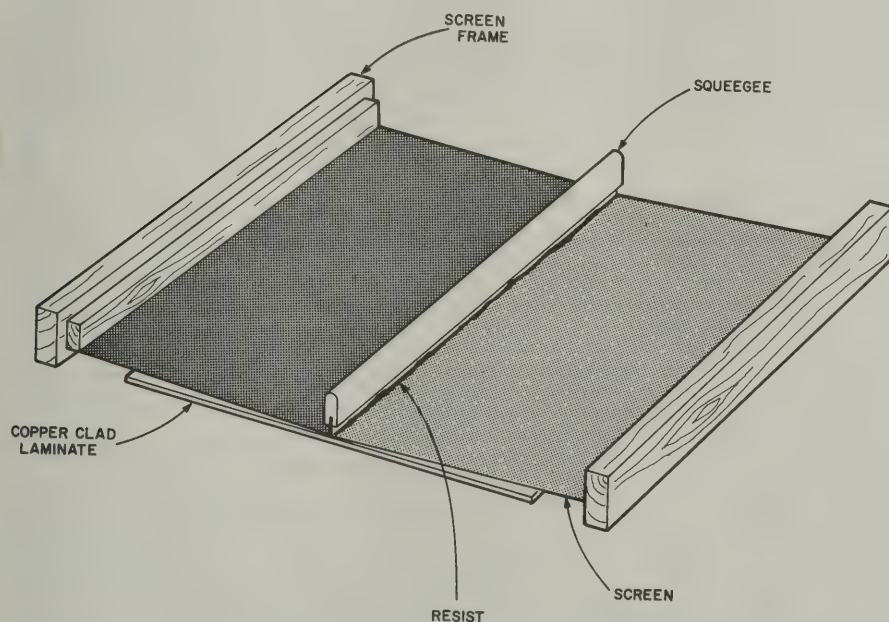


Figure 4

Figure 4 shows a typical setup for SILK SCREEN PRINTING by hand. Basically, the method is a stencil operation in which the resist material is transferred to the cladding through a stencil image. The stencil is firmly attached to the surface of a silk or nylon screen stretched drumhead-tight just above the copper-clad laminated board. Resist material is forced through the screen and the open areas of the stencil and onto the cladding by the wiping motion of a rubber squeegee.

Screen printing is adaptable as a low-cost print-and-etch method where very fine detail and definition is not required. This is, in general, true of all the print-and-etch methods. For very fine detail and high quantity production, another method, called the **PHOTO-ETCH** method, is used. Since the photo-etch method is widely used, the technique will be described in more detail than the simpler print-and-etch methods.

The Photo-etch Method

There are several steps which must be followed in the preparation of a printed circuit board using the photo-etch method. The basic steps are: (1) cleaning, (2) resist application, (3) prebake, (4) print (exposure), (5) development, (6) post-bake, (7) etching, and (8) resist removal. These steps describe the actual circuit board production. There are, however, other preparatory steps such as artwork preparation, and photography where the circuit pattern is photographed. The photograph is then used to transfer the image to the circuit board. The artwork for photographic copy is prepared oversized. Often this is two, four, or eight times greater than the final size of the printed circuit board. The artwork is then reduced to the proper size by photographing the image on film. This will be discussed in the following step-by-step procedure.

The artwork is produced, or "laid up", oversized so that greater detail and definition can be achieved. The accuracy of final images is generally required to be within standard tolerances of ± 0.005 inches. Thus, a large artwork drawing, which can be photographically reduced later, is easier to work with than a small drawing. Before the copper-clad laminate cleaning and resist application can begin, the master artwork must be made and corrected. The artwork is mounted on a camera copy board. Factors such as the type of camera lens and lens-to-copy distance then determine the amount of reduction in the image size when the photo is taken. A typical photo reduction setup is shown in Figure 5. When the picture is taken, a PHOTO NEGATIVE is obtained which is used to print the circuit board. The negative is inspected closely and retouched since small film blemishes would be reproduced in the copper cladding during printing. Faults in the negative produce a serious type of circuit board imperfection, called PINHOLING. Pinholing is caused by dark spots in the transparent (circuit pattern) portion of the negative. When the negative has been inspected and corrected, actual fabrication, according to the 8 steps previously mentioned, may proceed.

Copper-clad laminates require some type of cleaning in order to assure that the resist will not lift or peel. Mechanical cleaning such as sanding or other abrasive methods are often employed. A very smooth or mirror-bright surface is not desirable since the resist material will not adhere properly to smooth surfaces. Chemical cleaning is often more effective than mechanical cleaning. Typical chemical cleaners include trisodium phosphate, sodium

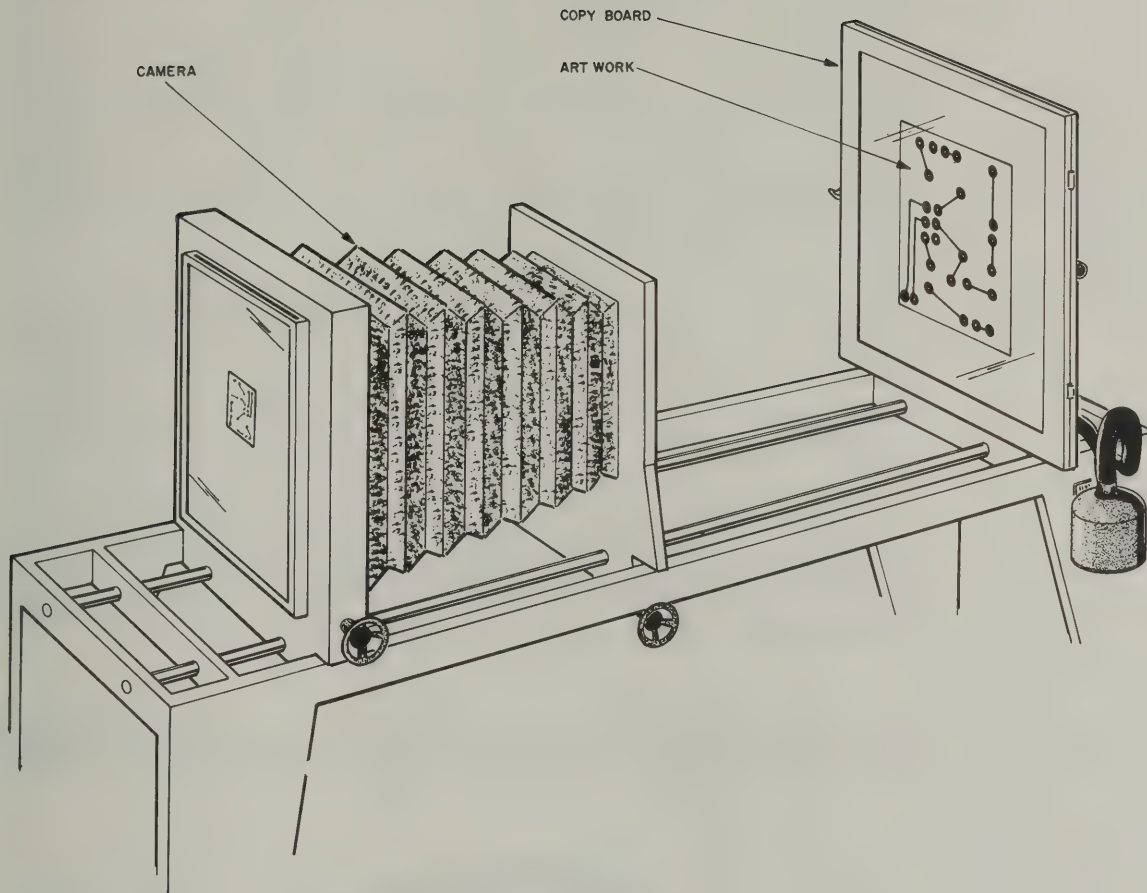


Figure 5

carbonate, and sodium hydroxide. Acid dipping and water rinses then follow to neutralize the alkali used in cleaning.

When the copper-clad laminate is clean, the photoresist is applied. Photo-sensitive resists are coatings which, when exposed to light of the proper wavelength, are chemically changed in their solubility to certain solvents (developers). There are two types of photoresists: negative acting and positive acting. Negative-acting resist is soluble in a developer prior to exposure to light. After exposure, the resist becomes polymerized, or hardened, and is insoluble in the same developer. Since a photo negative is a film with light and dark areas, a copper-clad laminate, previously coated with resist, will "print" on the resist just as on standard photographic paper. Thus, the dark areas on a photo negative block out the light to the circuit board when placed in a contact printer. The light area, which is the circuit pattern, allows the light to reach the board, and the resist becomes hardened

in these areas. When the exposure is completed, portions of the resist on the circuit board are soluble in a solvent developer, while other portions are insoluble in the developer.

Positive-acting resists are also used in photo-etch methods. This type of resist is initially insoluble in a solvent developer. Exposure to light, however, changes the solubility of the resist, and makes it soluble in the developer. This means that a PHOTO POSITIVE of the artwork is used rather than a negative, if used in the same system as the one described.

When the resist has been selected, the resist material is applied to the copper-clad board in a darkened room in preparation for exposure printing. The resist may be applied with a brush, roller or spray gun. The resist itself is a clear-to-amber-colored resin, resembling and handling much the same as varnish or lacquer.

A bake cycle prior to printing is usually used for photoresist coatings. This is the PRE-BAKE cycle. The board is first allowed to dry at room temperature. Pre-baking requires temperatures from 80° to 120°C (176° to 248°F), depending upon the type and thickness of the resist.

After the copper-clad laminate has been coated with photoresist and pre-baked, the board can be printed (exposed). In the case of negative-acting

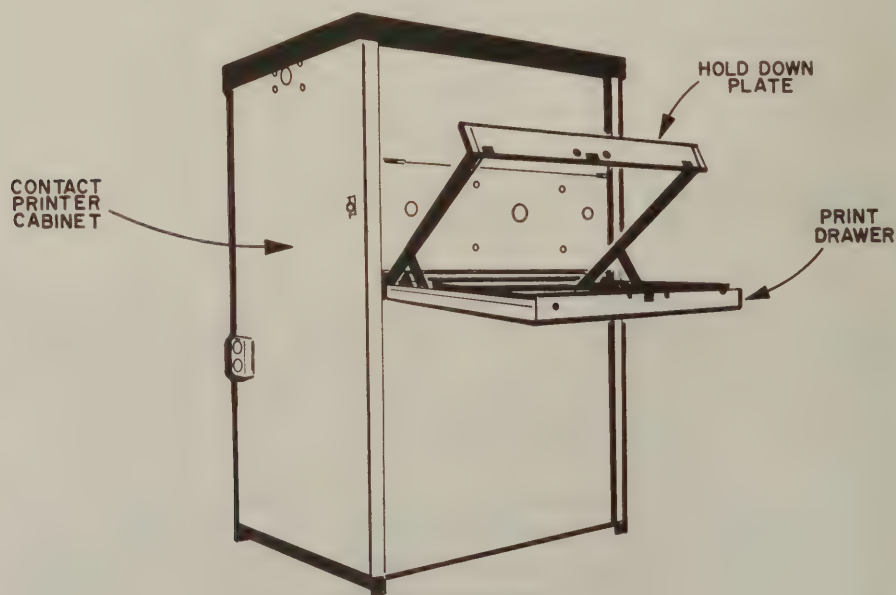


Figure 6

resists, the exposure causes hardening of the EXPOSED area, making it insoluble in the solvent developer. The board is most often CONTACT PRINTED. A contact print operation, as shown in Figure 6, is made directly onto the cladding. Relatively high-intensity light sources and long exposure times are used for exposing the resist. Generally used are carbon-arc, mercury-vapor, and, in some cases, tungsten lamps.

The developing operation consists of the removal of the photoresist material made soluble by exposure to light during the print operation. Types of solvents used include trichloroethelene and xylene for negative-acting resists. The time required for development varies with the type and thickness of the resist. Typical times are in the range of 2 to 3 minutes.

The circuit board may, for certain applications, be baked again after development to improve the chemical resistance of the image during the etch cycle. This is the POSTBAKE operation. This step is not always required. However, it is useful to remember that the postbake is not the same as the pre-bake cycle. The two steps are used for different reasons, and of the two, the pre-bake is the most important. The postbake is used to drive off any solvent which may have begun to penetrate the hardened resist areas.

The final critical step in the chemical processing of a copper-clad circuit board is the etch cycle. There are a variety of etch techniques. The simplest

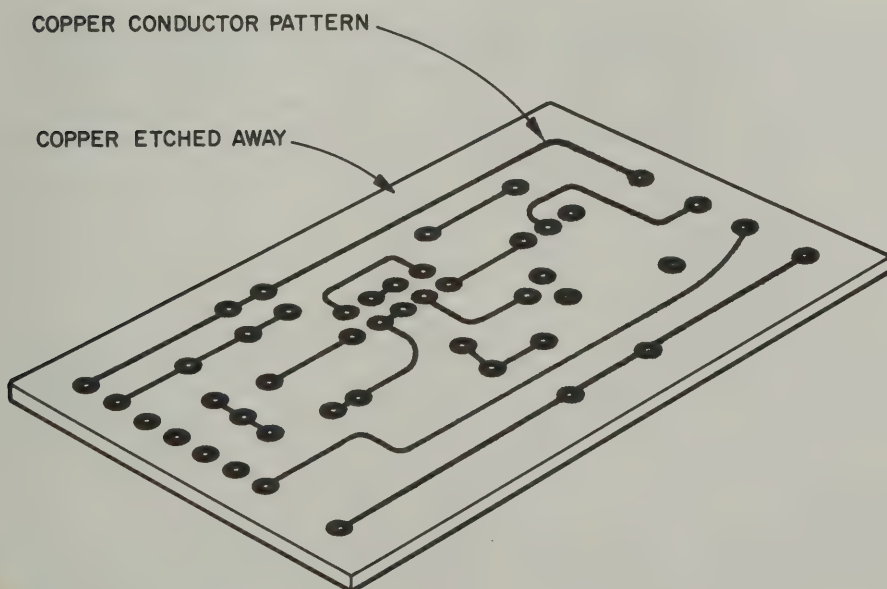


Figure 7

of these is the deep-tank or immersion etching process. With this method, the circuit boards are merely immersed in the etching solution until the etching is complete. Immersion etch techniques require long process times, and are suitable for laboratory work, but are rarely used in large-quantity production. Heating the etch solution and agitation of some form speeds up the action.

As has been repeatedly pointed out, the techniques and material used in printed circuit board fabrication vary widely. This is also true of etching solutions. The most widely used etchant for copper-clad laminates in print-and-etch processes is FERRIC CHLORIDE (FeCl_3). The composition of the etchant is essentially ferric chloride and water. The chemical reactions and theory of etching solutions is discussed in greater detail in the appendix. The etching solution oxidizes the copper which has been exposed by removal of the soluble resist coating. The etchant removes the exposed areas of copper, leaving the base laminate underneath. The copper under the hardened resist, shown as the dark area in the finished board of Figure 7, is not removed during the etch cycle.

When the board is processed completely, the circuit image remains as patterns of copper on the laminated board. The last remaining step is to remove the hardened resist remaining on the circuit-pattern copper. The resist must be removed because the resist would interfere with soldering and other operations performed on the board. Photoresist material on the image is removed using commercial strippers containing methylene chloride or other solvents. Scrubbing the surface with additional amounts of developer also removes thin coatings of resist.

Photo-etched circuit boards often use organic resists in procedures similar to those just described. These techniques produce satisfactory results for ordinary, non-precision type circuits. The procedures are also typical of laboratory prototype methods for single-board or small-quantity production.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

POINT-TO-POINT WIRING

COPPER FOIL TECHNIQUES

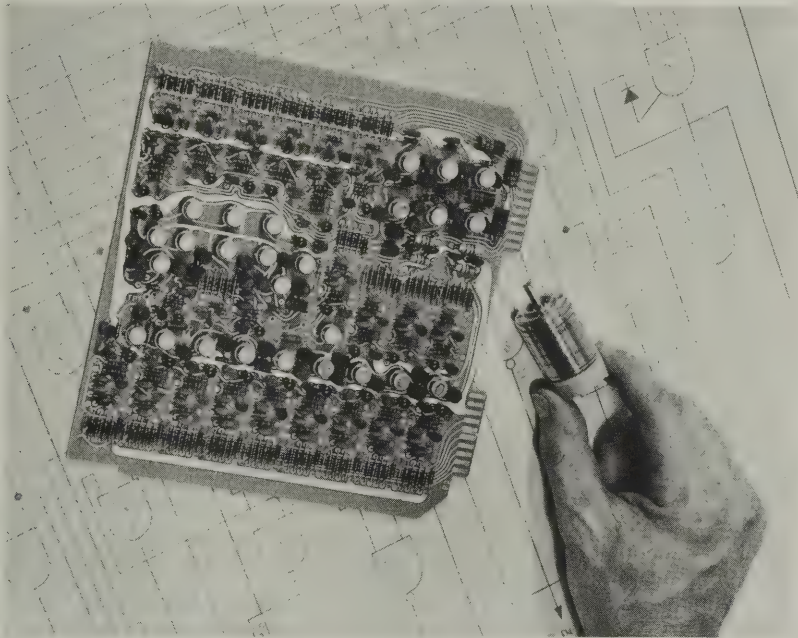
THE SILK SCREEN METHOD

THE PHOTO-ETCH METHOD

1. What two functions are served by using printed circuits?
2. The substrate used in a printed circuit board is called the _____.
3. In simple manual print-and-etch methods, the MATERIAL which is painted or taped directly onto the copper cladding is called the _____.
4. A faster print-and-etch method for producing larger numbers of printed circuit boards is to use a _____.
5. The ARTWORK used in photo-etch methods is the same size (1 : 1 ratio) as the final printed circuit. True or False?
6. Negative-acting resist is soluble in developer PRIOR to exposure to light. True or False?
7. The most widely used etchant for copper-clad laminates in print-and-etch processes is _____.
8. The etching solution oxidizes the copper under the HARDENED resist. True or False?

CIRCUIT LAYOUT

A printed circuit board is developed by determining the **CIRCUIT LAYOUT**. This is a procedure by which the final conductor routing and circuit configuration is produced. Ideally, the final circuit layout should be the



One of the useful properties of printed circuit boards is high component density packaging.

Courtesy North Atlantic Industries, Inc.

simplest physical representation of the circuit schematic. Simple circuits can be prepared for either single or double-sided copper-clad laminates. Another type of circuit layout uses multilayer boards. These are complex and are used in special applications. In this lesson we will be concerned mainly with single and double-sided circuit boards. When single-sided boards are used, the circuit layout is on only one surface or plane, and therefore any junction of two or more conductors establishes an electrical connection. This may be desirable only if an electrical connection is indicated in the schematic.

A schematic representation uses a dot symbol to indicate an electrical connection where two lines cross. If the dot is omitted, it is understood that there is no electrical connection at that point. This system, however, cannot be used in the circuit board layout. The use or omission of a dot at the

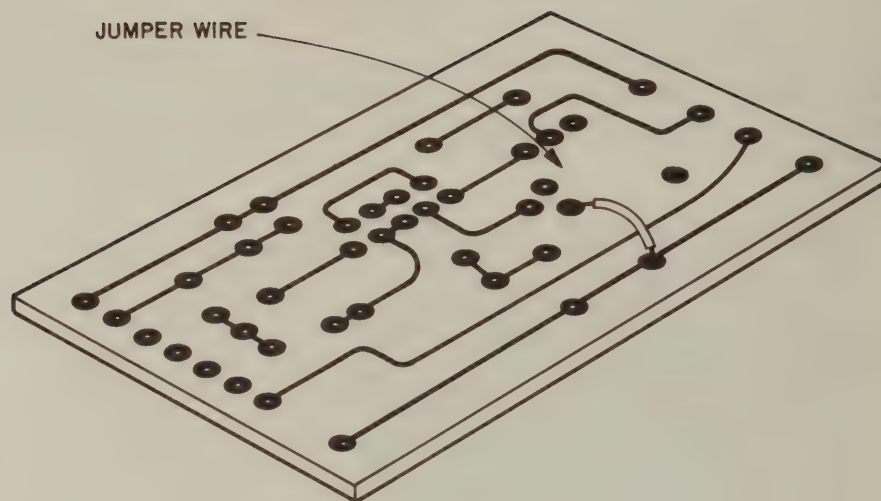


Figure 8

junction of conductors on a circuit board would be meaningless since the crossing of conductors on the board makes automatic connections. Because the conductors on a circuit board cannot cross in this manner, considerable planning is necessary to avoid the problem of conductor **CROSSEVERS**. Jumper wires may be used, as shown in Figure 8, where necessary, but these crossovers should be held to a minimum. Another method for minimizing the crossover problem is to use the components, especially the longer ones, to bridge the conductors which would otherwise cross. These problems and considerations are part of the layout procedure. Where the layout is too complex for adequately minimizing crossovers, a double-sided board may be used. A double-sided board permits the use of a second conductor pattern (one on each side of the board) and thus provides a means for crossing conductors where needed. Figure 9 shows a printed conductor acting as a crossover on a double-sided board.

A convenient aid to the layout procedure is to use a drawing of the components in a trial component layout. The circuit designer then may make free-hand point-to-point wiring diagrams while maintaining the sense and accuracy of the circuit schematic. Component rearrangement and repeated wiring diagrams are continued. Each successive redrawing results in a simpler circuit layout. Often two or three such initial steps are used, followed by two or three more carefully analyzed designs. Each step is merely a stepping stone toward the finalized design. When making the final design, or artwork, high quality drafting techniques must be used. If the artwork has unsharp edges for copper conductors, the resultant pc boards are much more likely to produce arcing if high voltage testing or use is intended.

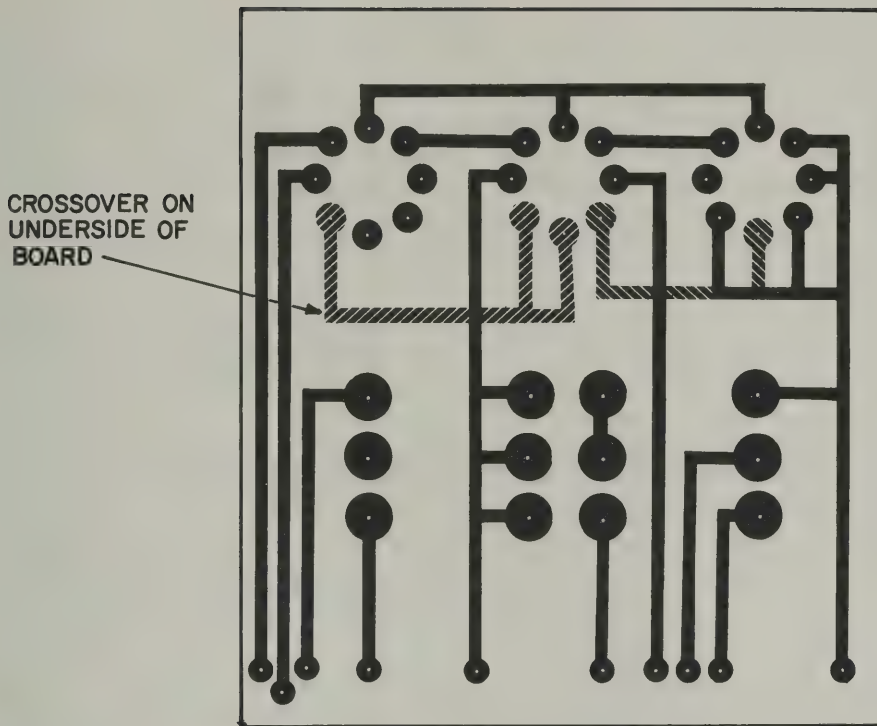


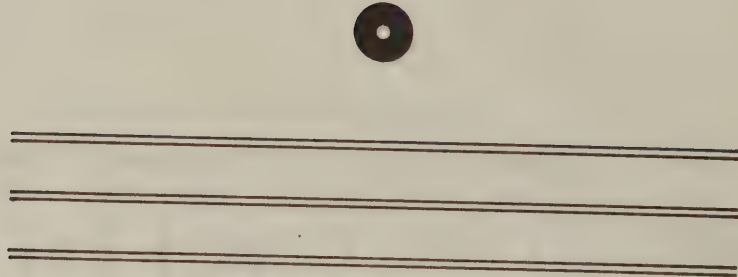
Figure 9

The electrical characteristics of component and conductor placement also must be considered carefully since the distances between components and conductors are small. Effects of mutual coupling, inductive and capacitive, must be considered for some types of circuits. The insulation resistance and dielectric properties of the substrate require consideration for high frequency and high impedance circuits.

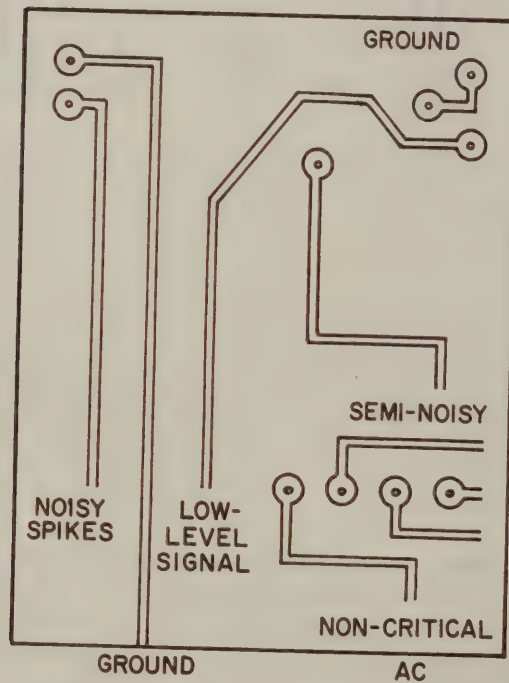
There are other techniques within the layout procedure which influence good design in printed circuitry. These considerations include mechanical placement of fixed components such as trimmer potentiometers, switches, and test jacks. Positions of these components should be such that accessibility is not impaired. However, copper paths to these components should be reasonably short to avoid undesirable electrical faults such as noise pickup or shunting of signals by stray capacitances.

CONDUCTOR ROUTING

The same principles that make a good layout for point-to-point wiring apply also to printed circuit wiring. Methods are required to guard against undesirable interference by using either adequate separation or shielding. Although the principles are the same, the physical layout is different. Since



PARALLEL-RUN CONDUCTOR CONFIGURATION
A



CONDUCTOR ROUTES UTILIZE
SPACE AVAILABILITY

B

Figure 10

there is only one plane, more circuit planning is required to assure acceptable circuit performance. Conventional point-to-point hookup wire is insulated and may be routed on intersecting paths without regard for the intersections (except where mutual inductance is a factor). Printed circuit wiring, however,

establishes an electrical connection at such junctions. To overcome this, conductor routing requires some thoughtful planning. Many times, component leads will be longer than desired due to the attempt to span several conductors on the reverse side of the board. Also, some conductors may be longer than desired due to complicated routing requirements. These are instances, however, where a jumper wire or crossover may be used to advantage. Additionally, the conductors should utilize as much of the available space as possible. To do this, each conductor run should be equally spaced from the adjacent conductors. This means that, ideally, the conductor route should split the difference at every conductor turn, and have equal spacing over parallel runs, as shown in Figure 10A. Another consideration in determining conductor routing is the use of the available space on a circuit board. If a relatively large board is to be used but has only a few conductors on it, the full area of the board should be used. Needless crowding should be avoided. Figure 10B shows a functional layout which uses the board area to its best advantage. Notice that special consideration such as noise and voltage levels are handled by spacing and separation. This design makes the best use of the available space, while giving the most consideration to the circuit characteristics.

Another consideration in conductor routing is that of DIMENSIONS. Narrow conductors and close spacing increase the cost of a printed circuit board. However, the actual dimensions used in any given application are determined by a number of other design considerations. One of these considerations is the current requirement in the printed conductors. Relatively large currents require wider conductors and/or thicker cladding material in order to minimize power losses and heating. Another governing factor is accuracy achievable by various handling or processing machines which may be in use. Hole drilling or component insertion machines may be limited in placement tolerance characteristics. When this is the case, the conductor widths and spacings may have to be larger than would be required if current-carrying capacity were the only factor. For many applications, 1/32 inch spacing between conductors has been found satisfactory. Often a 1/32 inch conductor width is also used.

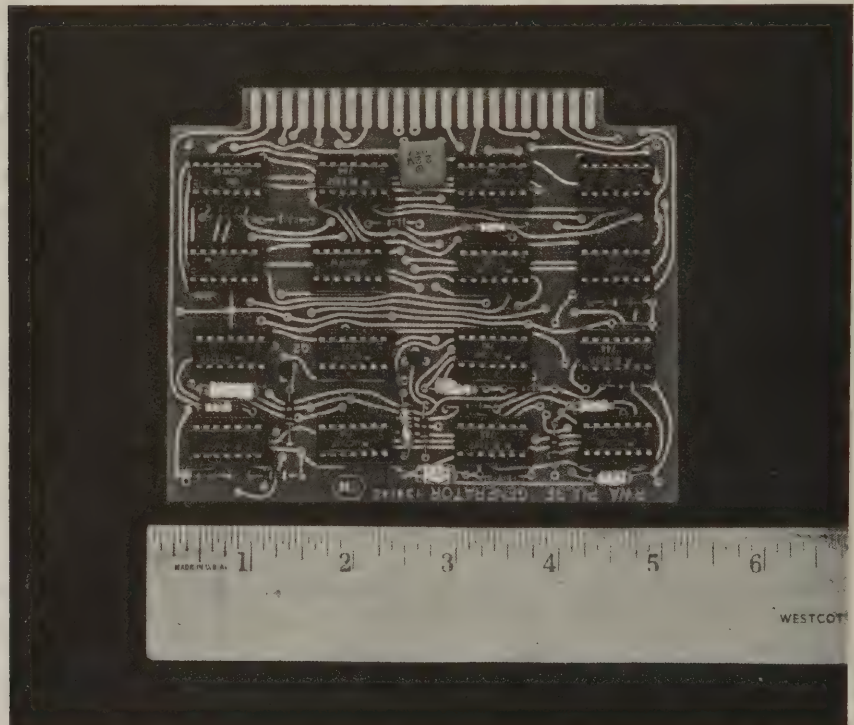
Conductors that are overly wide can cause soldering problems because the copper may act as a heat sink. This may necessitate the use of a hotter soldering iron when attaching the components. As we shall see in a later section, excessive heat can cause the conductor to separate from the board. Thus, an untrained or inexperienced assembler may have difficulty in properly soldering some printed circuit boards. Where large areas of copper are required, such areas are often shaped into a crosshatch pattern to avoid excessive heat transfer or bubbling up, or lifting, of the copper if there still would be excessive heat.

Extremely close spacing between conductors may encourage solder BRIDGING. This occurs when excessive amounts of solder are used when attaching

the components, and very close spacing increases the possibility of bridging, or solder spilling over from one conductor to an adjacent conductor during component assembly. This may be a very small amount of solder and may go undetected by the assembler. Generally, however, even a small amount of solder is sufficient to cause a short between the conductors where none should exist.

The ~~copper~~ thickness is often referred to as "one ounce, two ounce," etc. Two ounce copper means a copper cladding on the laminate of a thickness such that one square foot of the plating weighs two ounces. The copper that is used has a thickness of approximately .0014 inches per ounce. Thus, two ounce copper is .0028 inches thick. A .032 inch wide two ounce copper conductor can safely carry a current of two amperes.

In general, printed circuit board conductors should be laid out so that the finished product provides all the required circuit interconnections in the simplest possible form. This requires the consideration of still another dimension in conductor layout: conductor SHAPE. Printed circuitry uses



Conductor shape is often an important factor in the overall design of printed circuit boards. The smooth, flowing design shown here makes the best use of the available space.

Courtesy E-H Research Laboratories, Inc.

either curved, flowing conductor shape, or straight, angled shape. Both types have specific advantages and disadvantages, but either may be used. An advantage of the smooth line technique is that the conductor may arc or sweep directly to the required connection point. This often reduces overall conductor length. A disadvantage of the smooth line technique is that the artwork is more difficult to prepare. An advantage of the straight line technique, however, is that the artwork is relatively simple to prepare. Flat and straight tapes can be used in the artwork, and do not require specialized drawing skills. A disadvantage of the straight line method is the increased probability of electrical potential concentrations which encourage arcing. Manufactured boards are often coated with a solder resist which prevents undesired areas of the copper from being soldered.

PLATING

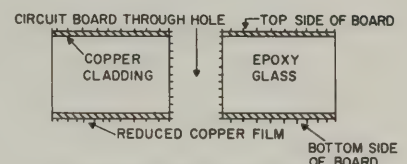
Circuit board plating refers to the techniques for obtaining metallic BUILD-UP in the drilled holes of a copper-clad laminated circuit board, and also on the circuit pattern. Metal plating on the walls of the hole is needed to provide a conductive path between circuits on opposite sides of double-sided, and multilayer circuit boards. The general process is called **THROUGH-HOLE PLATING**.

Through-hole plating is part of a high-quality circuit fabrication process, and involves some additional steps. These steps are concerned mainly with the plating of metal on the nonconductive hole walls in a pre-drilled, laminated board. Several methods exist for plating on nonconductive surfaces. These include the use of graphite, conductive lacquers, metal powders, chemical reduction, silver spray, and vapor plating. The technique most often used in through-hole plating is the CHEMICAL REDUCTION OF COPPER.

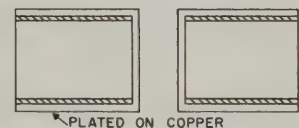
The copper reduction process involves **SENSITIZING** the nonconductive hole surface. The sensitizing process deposits tin (SN^{+2}) ions on the surface of the laminate. The nonconductive surface will then accept a thin film of a precious metal such as palladium or gold. The next step is the plating of a thin copper layer over the palladium or gold by copper reduction in a copper salt solution. When the plating process is complete, the conductive coating in the holes "bridges" the copper cladding on both sides of the board.

There are two general techniques which use the copper reduction process to plate printed circuit through holes. These are PANEL PLATING and PATTERN PLATING. The series of diagrams in Figure 11 illustrates the steps used in panel plating. In this procedure, the entire copper cladding is sensitized and copper plated using the copper reduction process as described previously. After a film of copper has been deposited in the drilled hole and over the copper cladding (Figure 11A), a heavy layer of copper is electro-

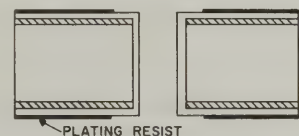
PANEL PLATED



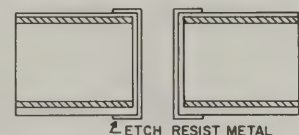
A SENSITIZED AND METALLIZED



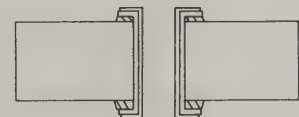
B COPPER PLATING



C NEGATIVE IMAGE PRINTED



D ETCH RESIST METAL PLATED AND STRIPPED



E ETCHED-FINISHED CIRCUIT

Figure 11

plated onto the entire surface of the board (Figure 11B). The thin copper film provides the base to which the electroplated copper layer can adhere. The plated panel is coated with photoresist and printed with the circuit pattern using the normal photo-etch contact print method. This step is shown in Figure 11C. When a photonegative is used in printing, the photoresist used is **POSITIVE-ACTING**, and thus is made soluble by exposure to light. After exposure, the photoresist is de-polymerized in the through-hole and on the circuit pattern. The resist in the through-hole and on the circuit pattern is exposed to the light, and becomes soluble. When the image is developed, the resist is **REMOVED** from the through-hole and circuit pattern. The board is then ready for another plating step. The exposed photoresist is able to withstand, or "resist" this second plating step which follows development of the image. This plating covers only the through holes and the circuit pattern, and does not adhere to the hardened resist. This plating is a **LEAD-TIN**, or **SOLDER** plate. The lead-tin or solder plate adheres only to the hole walls and circuit pattern where the photoresist has been removed as shown in Figure 11D. This plating is called the **ETCH-RESIST METAL** because the solder plate is not affected by the final etch cycle. When the solder plating has been applied, the remaining photo-resist is removed. This process is called **STRIPPING**, which removes the hardened resist and exposes the copper on the unwanted areas. The final step, shown in Figure 11E, is the etch cycle which removes all the unwanted copper including the exposed copper plating, reduced copper film, and the original copper cladding. The part which remains is the solder-plated through-hole and the circuit pattern on **BOTH** sides of the board.

Pattern plating accomplishes the same end result as panel plating, but requires less copper and is somewhat more economical. Pattern plating deposits the plated layers only on the hole walls and the circuit pattern. The steps for pattern plating are shown in Figure 12. Note that step A is the same for both procedures. Step B in Figure 12 is similar to step C in Figure 11. This is the photoresist application and printing with the circuit pattern. Note that in pattern plating, the electroplated copper coat is not applied until **AFTER** the photoresist has been applied. The photoresist (also called the plating resist) serves as the resist for both the copper plating and the solder plating. The two plating steps are performed in succession. With the through-hole finally plated with solder, the hardened photoresist is stripped as shown in Figure 12D. The final step is the etch cycle and is the same for both procedures.

CIRCUIT TRACING PRINTED CIRCUIT BOARDS

Once a printed circuit board has been designed and fabricated, the assembled circuit must be tested to detect flaws. These may be as a result of the circuit design, or due to fabrication errors. A great number of processing steps are used in printed circuit fabrication, and flaws can occur in any of these steps. Printed circuit board testing is an important part of the manufacturing process.

PATTERN PLATED

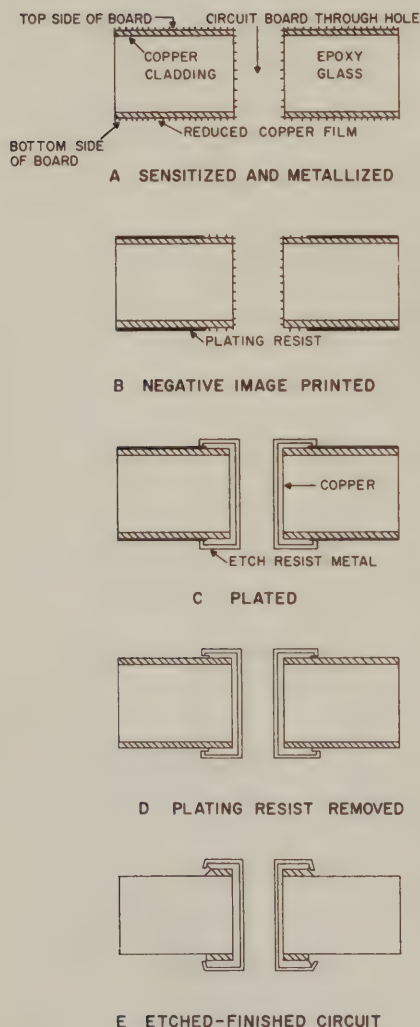


Figure 12

A completed printed circuit board is best tested for faults by thoroughly checking it under operational conditions. This includes the required power supply connections, signals, and circuit load conditions.

There are also several procedures which, if observed during the circuit layout stage, greatly simplify the testing operations. One of these is the use of test points. Test points should be included in the original circuit design, and provisions made for them during the artwork stage. The test points should be located in accessible places, preferably near the edges of the board.

Another consideration is to avoid connecting components in parallel on the same printed circuit. An example of this is shown in Figure 13. The parallel component layout in Figure 13A cannot be tested accurately due to the parallel circuit path through the transistors. Note that if either transistor were faulty, the other one could provide an output during the test. One solution to this problem is to connect the circuits together externally off the board as shown in Figure 13B. This permits at least one of the transistors to be tested separately. Then with one circuit "known", it can be determined whether or not the other one is working properly.

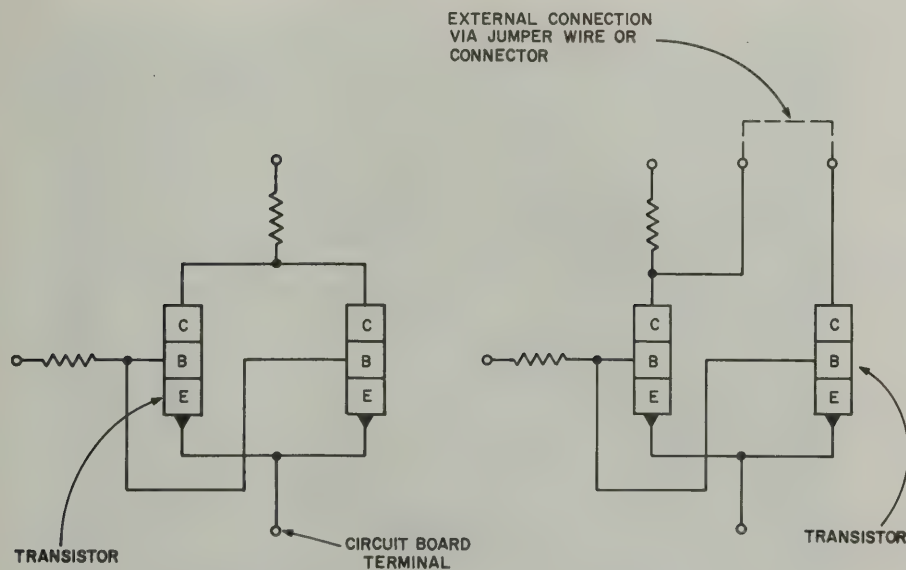


Figure 13A

Figure 13B

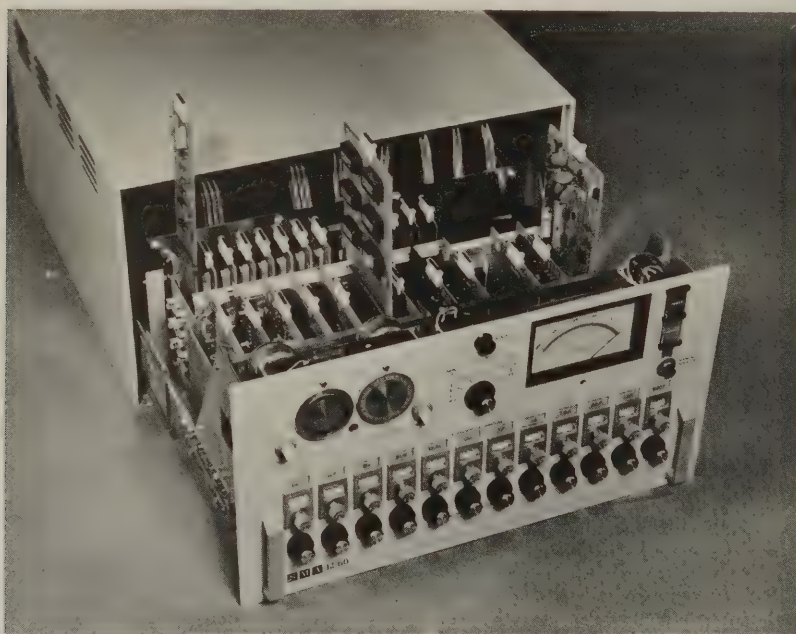
Most test procedures for printed circuits use conventional test equipment to provide input signals and voltages to the circuit. These instruments include sinewave oscillators, ramp generators, and pulse generators. Voltage levels

and passive input conditions may be simulated using power supplies and voltage dividers. Outputs are tested by simulating the circuit output loading conditions. In some cases, actual operating circuits must be used to test loading. Such load simulators are thoroughly pre-tested and are known to be good.

The actual circuit board testing may be performed manually or by automatic circuit card testers. Manual handling and probe contacting is the simplest method, and is suitable for low quantity work. A special mating test socket is often used to simultaneously hold the board and apply the test voltages. An operator connects the board to the test station, and performs the required tests by manually probing the circuit. Automatic testing requires the use of sophisticated equipment. This equipment automatically HANDLES, or positions, the board, and then applies the test voltages. Fully automatic equipment may display the test readings for operator interpretation, or the results may be machine-analyzed.

REPAIRING PRINTED CIRCUIT BOARDS

Printed circuit boards are convenient to manufacture and use, but require some special considerations in repair and maintenance. For one thing, al-



Provision is often made for vertical plug-in, plug-out interconnection within the equipment. These racks, also called drawers, permit easy removal for inspection and repair.

Courtesy Datascan, Incorporated

though the copper conductors are quite firmly attached to the substrate, they can be pulled loose, cracked, or become damaged in other ways. Thus, there are two aspects to the problem of printed circuit board repair. One of these is the replacement of faulty components. The other is the repair of the board itself.

Replacing defective components on a printed circuit board requires special care because excessive handling can damage the board. Prolonged heating or use of a soldering iron which is too hot can cause the conductor to lift. This situation is very difficult, if not impossible, to repair satisfactorily. Also, the components on many factory-assembled printed circuit boards are held in place mechanically by a band, or hook, in the leads on the underside of

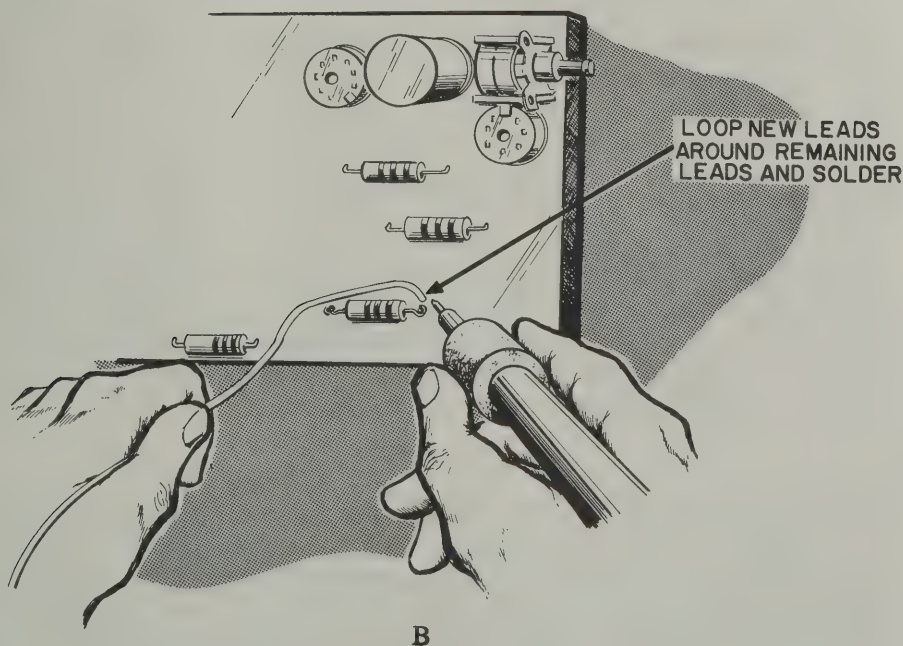
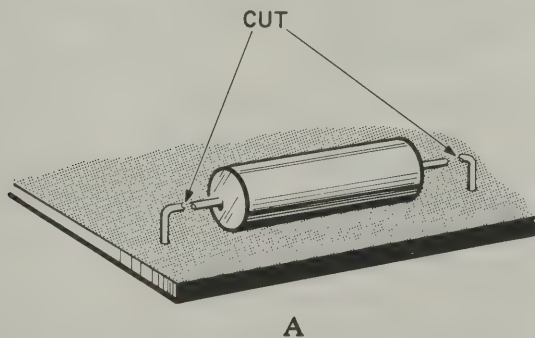


Figure 14

the board. For these reasons, one should not attempt to completely remove a component by de-soldering with a soldering iron. Instead, the technique shown in Figure 14 is preferred.

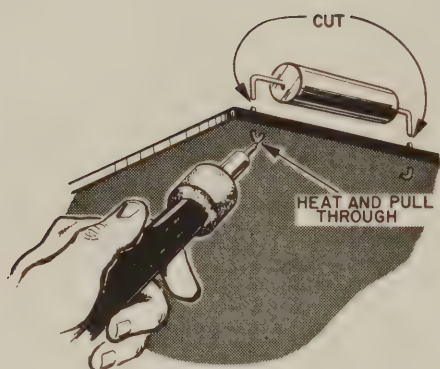


Figure 15

A component such as a resistor or capacitor is removed by cutting the tubular part out while leaving the ends of the original leads soldered in the board. The remaining leads provide a type of terminal to which the leads of the new part can be attached. A common technique for attaching the new component is to form a small loop in each of the new component leads, as shown in Figure 14B. The looped ends can be slipped over the old leads and soldered in place. The loop provides a good mechanical connection which holds the component while soldering, and also improves the quality of the electrical connection. The loop can be formed by wrapping the lead around a stiff wire or another component lead.

Occasionally a component must be removed from the board completely. When this is the case, the component can still be removed without damaging the board by cutting the component out as before. The lead should be cut closer to the board, however, so that only a short vertical stub is left. Then, as shown in Figure 15, the solder around the remaining lead will melt easily and can be quickly brushed away. The exposed lead can then be removed from the **UNDERSIDE** by pulling the straight portion of the lead through the hole, while being careful to not pull the copper conductor away from the board if the lead should catch it while being removed.

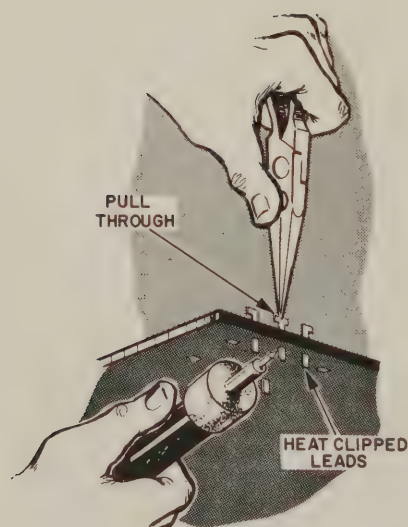


Figure 16

Other components such as r-f coils, transformers, tube sockets, and potentiometers are more difficult to remove. The correct procedure is shown in Figure 16. These components often use short lugs which mount directly on the board. The lugs should be cut off from the underside, assuming that the component is not to be repaired. If the component is to be repaired, the lugs will have to be straightened, and the component removed more carefully. Due to the rigidity of the component, the lugs generally must all be removed simultaneously. This means that the solder around each of the lugs must be melted and brushed away until the lugs are free. The excess solder can also be removed with a solder sipper. This should be done as quickly as possible to avoid overheating the copper conductor. When the lugs are loose, the component should separate easily from the board.

There are several aspects to the problem of repairing a printed circuit board. The circuit board is, itself, as much an electronic component as is a capacitor or transistor. This is because the board, while an integral part of the whole circuit, actually develops its own unique set of characteristics. Thus, the circuit board is subject to failure and requires special know-how to repair and maintain.

A common repair problem with printed circuits is in locating cracks in the conductors. Cracks often develop which are so small that they are difficult to locate visually. These cracks, although small, can cause open or intermittent circuit conditions. The fault can usually be located by placing the circuit board on a light table (a glass top table with a light behind it), or by simply holding the board up to a strong light. Since most substrates are translucent, the fault shows up as a thin line of light across the conductor. When the crack has been located, a suitable repair method is selected.

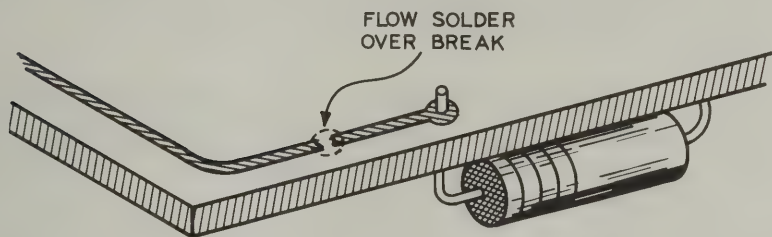


Figure 17

If the crack is very small and the conductor has not lifted from the substrate, a small amount of solder flowed over the crack may be sufficient as shown in Figure 17. Slightly larger cracks may be repaired by using a short length of tinned wire as shown in Figure 18. Do not attempt to use stranded wire. In these short lengths, stranded wire has a tendency to separate and it will be very difficult to keep the strands from shorting to another conductor.

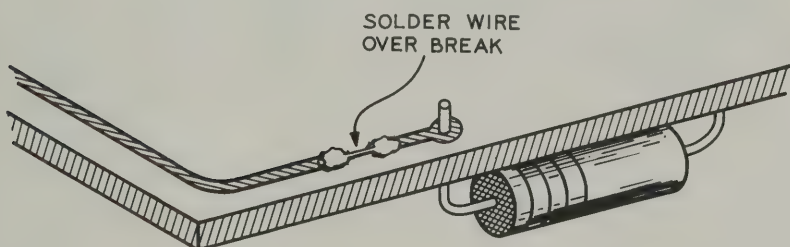


Figure 18

In cases where the conductor has been damaged more heavily, these methods may not work. In fact, if the conductor has lifted at the break, the heat from the soldering iron may only make the problem worse. Once the conductor has lifted from the substrate, the loose section must be removed and a jumper

wire soldered in its place. This is shown in Figure 19. The procedure is as follows:

1. Remove all solder from the conductor in the break area by heating the conductor and brushing the solder away, or by use of a solder sipper. If the damaged section occurs near a terminal pad, the through hole should be cleaned of solder.
2. Using a sharp blade, cut across the conductor slightly behind the lifted section. This should sever the conductor, allowing the damaged portion to be removed completely.
3. Solder an insulated jumper wire from the secure end of the conductor to the component lead in the through hole. A slight twist of the wire around the component lead will hold them together while soldering.

Another printed circuit board trouble which appears is leakage resistance. This is similar to the capacitor leakage resistance especially noticeable in electrolytics. Due to the close proximity of the conductors and the dielectric properties of the substrate, printed circuit boards can cause special problems in high-frequency work.

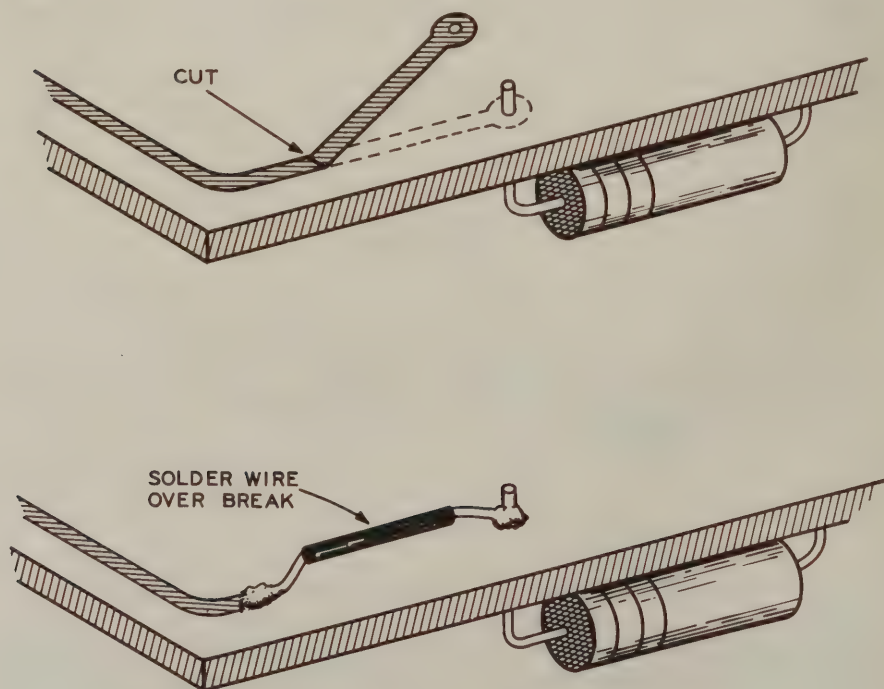


Figure 19

The following Practice Exercise questions cover the subjects which you have just studied. They are:

CIRCUIT LAYOUT

CONDUCTOR ROUTING

PLATING

CIRCUIT TRACING PRINTED CIRCUIT BOARDS

REPAIRING PRINTED CIRCUIT BOARDS

9. Printed circuit layouts are electrically the same as the associated circuit schematic. True or False?
10. The reason that a circuit board layout looks different from a schematic drawing is that there are no conductor _____.
11. In addition to the circuit pattern layout, the components must be placed in such a way as to avoid unwanted mutual _____.
12. Conductor widths should not be smaller than $\frac{1}{8}$ inch. True or False?
13. Plating refers to the process of building up a metallic plate in the drilled holes and on the circuit pattern of a high-quality printed circuit board. True or False?
14. The plating technique which deposits a copper layer over the entire copper-clad surface is called _____.
15. Pattern plating is less economical than panel plating. True or False?
16. When removing a tubular-component from a printed circuit board, the best procedure is to melt and brush the solder away from the pad until the component can be lifted from the board. True or False?
17. Leakage resistance in a printed circuit board can be caused by (a) an open resistor, (b) metallic dust between conductors.

Leakage resistance can be caused by microscopic metallic dust or other foreign material collecting between the conductors. This forms a high resistance path between components which should be electrically insulated from each other. It also may cause leakage paths for capacitors, thus making a capacitor appear faulty when, in fact, the circuit board is at fault. This condition can also occur from defects in the substrate or from resin buildups on the board.

If leakage resistance is suspected on a faulty board, it can usually be detected by using a meter capable of reading high values of resistance (up to 1,000 megohms). Once the faulty area has been located, thorough cleaning with solvent and a stiff brush may correct the problem. The brush should not be so stiff, however, that it scratches the copper. A toothbrush, or similar brush, is preferable.

Circuit boards also develop high-resistance connections where a low resistance connection (short) should exist. This can be due to a partially-cracked or corroded conductor. When the entire cross-sectional area of a conductor is not capable of carrying current, the conductor acts as though a section of high-resistance wire were inserted in the circuit. Again, this situation may be difficult to spot with the eye. A quick visual check may reveal an area which looks suspect, but the final check generally requires the use of an ohmmeter. By placing the probes on either end of the suspected conductor section, the meter should read zero ohms regardless of the component configuration on the board.

In some cases, a high resistance conductor will appear faulty only intermittently. This may be due to the stress changes on the board when the mounting lugs are tightened, or by inadvertent bending of the board. Since the conductors are affixed to the substrate, bending may stress the conductors and cause them to crack. These cracks can cause an open circuit condition when the board is stressed in one direction, but may perform properly when the board is stressed in the other direction. This condition is probably the most difficult of all circuit board problems to locate and repair.

An intermittent circuit condition may be found locally by gently tapping with an insulated tool on the conductors and soldered connections while the board is operating. By monitoring the output and noting any changes which may occur as a result of the tapping, the general area of the trouble can be found. If, upon closer inspection of the suspected area no cracks can be found, all the conductors in that area of the board can be tinned with solder, the connections reheated and the test repeated. By tinning as much of the copper surface of the board as possible, the fault may be corrected without actually locating an otherwise undetectable flaw.

After a circuit board has been repaired, the work should be thoroughly examined to see that excess solder does not remain between conductors where it can cause further trouble.

SUMMARY

The development of more sophisticated electronic equipment has required the use of smaller and more reliable components. Printed circuit techniques have played an important part in electronic miniaturization. A printed circuit board serves a dual function. It provides the mounting space for the components and also connects them together into a working circuit.

Although the electrical function of a printed circuit board is the same as could be obtained with point-to-point wiring, the physical layout is different. Simple printed circuit boards are produced on only one plane. The conductive pattern is etched in metal on a copper-clad laminate by print-and-etch or photo-etch methods.

The photo-etch method is capable of producing better registration (pattern accuracy) than print-and-etch methods. The steps required to produce a printed circuit board using the photo-etch method are:

1. cleaning
2. resist application
3. prebake
4. print (exposure)
5. development
6. post bake
7. etching,
8. resist removal

The image to be transferred to the board is made photographically. Therefore, a photo negative of the pattern must be produced just as a negative is required in order to produce an ordinary photograph. The artwork, or master pattern, is produced oversized using standard drafting techniques. A negative of the proper size is produced by photographing the artwork.

The copper cladding is thoroughly cleaned prior to application of a photo-resist coating. After a pre-bake, the board is exposed in a contact printer. The circuit pattern is developed and baked again, which prepares the board for the etch cycle. The final critical step in the chemical processing of a printed circuit board is the etching cycle. This step removes the unwanted copper from the board and produces the final copper pattern.

A circuit designer must produce the pattern which is transferred to the laminate while maintaining the sense and accuracy of the circuit. This is accomplished by producing a series of freehand drawings, each one becoming more refined and finalized. It is in the layout stage that various circuit design problems and conveniences must be considered. Also, the same conductor routing principles that make a good layout for point-to-point wiring apply to printed circuit wiring. Methods are required to guard against undesirable interference by using adequate separation and shielding. Conductor dimension and shape contribute to the considerations in producing a well-designed circuit.

Plating is a method for producing high-quality double-sided and multilayer circuit boards. The plating procedure deposits a layer of metal in the through hole of the board. In addition, the plating produces a high quality circuit pattern. The two general techniques used in plating are called panel plating and pattern plating. The final results are the same using either method. However, pattern plating is the most economical of the two.

When a circuit board has been produced, it must be tested to see that it performs as expected. This is most easily done by checking it under operational conditions. Where test points have been provided during the layout stage, various internal voltage and signal conditions can be checked. A variety of manual and automatic circuit card testers are in use to speed the testing of commercial circuit boards.

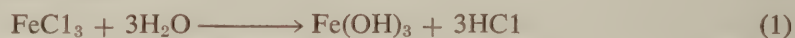
Printed circuit boards are subject to types of failure which are uniquely their own. Because of this, circuit boards can be considered another type of passive circuit component. Often a circuit may develop a problem which produces the symptoms of resistor, capacitor, or transistor malfunction when the problem actually results from a printed wiring malfunction. When a resistor or capacitor malfunctions and must be replaced, special care must be exercised in removing the component. When the board itself malfunctions, specialized repair techniques must be used. A common printed wiring failure is the development of small cracks in the conductors. These cracks may be repaired by flowing a small amount of solder over the crack, or by using jumper wires. Other troubles include leakage resistance and high resistance connections where a low resistance connection should exist. A circuit conductor may fail completely, or the failure may be intermittent in nature. Of the two, an intermittent failure is the most difficult to locate and repair.

APPENDIX

Copper Etchant Reactions

Ferric chloride (FeCl_3) is among the most widely used etchants for copper and copper alloys. A major reason for this is because FeCl_3 has a high tolerance for dissolved copper. The solution can be used repeatedly and has a long useful life.

The etchant is composed of FeCl_3 in water. Free acid is present due to the hydrolysis reaction

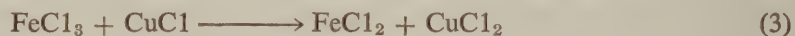


Additional amounts of HCl , hydrochloric acid, may be added to reduce the formation of insoluble precipitates of $\text{Fe}(\text{OH})_3$, ferric hydroxide.

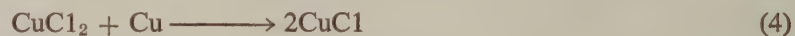
In the oxidation reaction between the ferric chloride and the copper surface, the ferric ion (Fe^{+++}) oxidizes copper to cuprous chloride and ferrous chloride.



The ferrous chloride is responsible for the greenish tint of the brown liquid. The cuprous chloride is further oxidized to cupric chloride (CuCl_2).



As the concentration of cupric chloride builds up in the etching solution, a partial breakdown reaction takes place. Cuprous chloride (CuCl) replaces cupric chloride (CuCl_2) in the presence of elemental copper (Cu).



Eventually, the major reaction in the etching of printed circuit boards takes place according to Equation 4.

Ferric chloride etching solutions vary from one manufacturer to another depending upon special additives used. However, the etch time for a printed circuit remains relatively constant up to 40 percent exhaustion (11 oz. copper/gal.). Although the etching time increases rapidly beyond 11 oz. copper per gallon, the solution still has etching strength. Because of this, a variety

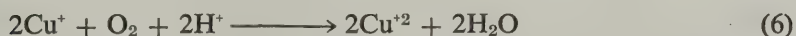
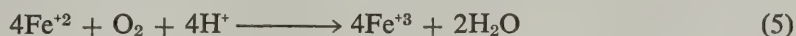
APPENDIX—(Continued)

of procedures have been developed to increase the usefulness of partially-depleted etching solutions.

1. Reduction of the etching time required of partially-depleted solutions by increasing the temperature of the solution. The temperature can be adjusted to obtain the desired etch time without damaging the resist material, or producing excessive fuming.

2. Reduction of the etching time required of partially-depleted solutions by the addition of HCl. When the copper concentration exceeds approximately 6 oz./gal., HCl reduces the formation $\text{Fe}(\text{OH})_3$ and enhances the etching reactions.

3. Reduction of the etching time required of partially depleted etching solutions by agitation. The time required to etch a given amount of copper in an 8 oz./gal. still bath is approximately one third greater than that required for a 16 oz./gal. of agitated solution at the same temperature. Etching time can be reduced still further by using air agitation in a spray, splash, or bubble operation. The etch rate is enhanced by the oxidation of Fe^{+2} and Cu^+ to Fe^{+3} and Cu^{+2} .



IMPORTANT DEFINITIONS

ARTWORK — The master pattern, or lay-up, of the final etched circuit which is produced by using conventional drafting techniques.

CLADDING — The metallic foil layer which is bonded to the substrate and becomes the conductive pattern after etching.

CROSSEOVERS — Jumper wires used on a printed circuit board where a conductor must cross one or more other conductors.

ETCHANT — The solution which attacks and removes selected portions of printed circuit cladding.

LAMINATE — The substrate of a printed circuit board commonly made of phenolic paper, epoxy paper, or epoxy glass.

MINIATURIZATION — The use of smaller electronics components such as miniature tubes and solid state devices to produce smaller and more reliable equipment.

OXIDATION — The raising of an ion or a neutral atom to a more highly charged (positive) state. Oxidation occurs when an ion goes from an anion to a neutral atom, a neutral atom to a cation, or a cation to a more highly charged cation through the LOSS of electrons. (e.g. ferrous⁺⁺ to ferric⁺⁺⁺.)

PADS — Circular enlargements in a conductor suitable for mounting and soldering a component lead to the conductor at that point.

PANEL PLATING — A plating technique in which a heavy electroplated layer of copper is applied to the board PRIOR to the printing cycle.

PATTERN PLATING — A plating technique in which a heavy electroplated layer of copper is applied to the board AFTER the printing cycle.

PHOTO-ETCH — A photographic process for transferring a printed circuit pattern image to a copper-clad substrate.

PINHOLING — A circuit board fault caused by imperfections or foreign substances on the photo negative.

POINT-TO-POINT WIRING — The conventional hand-wired chassis inter-connection method.

PRINT-AND-ETCH — All methods of producing printed circuits other than the photo-etch method.

IMPORTANT DEFINITIONS—(Continued)

REDUCTION — The lowering of an ion or a neutral atom to a less highly charged (positive) state. Reduction occurs when an ion goes from a cation to a less highly charged cation, a cation to a neutral atom, or a neutral atom to an anion through the **GAINING** of electrons. A typical reduction reaction occurs in the electroplating process.

RESIST — A material which is resistant to the etching acid used in making printed circuit boards.

THROUGH-HOLE PLATING — A process for building up a layer of conductive material on the walls of the drilled holes in a printed circuit board.

PRACTICE EXERCISE SOLUTIONS

1. Printed circuits serve to support the components and provide the conductive paths between the components.
2. laminate.
3. resist.
4. silk screen.
5. False — The artwork is produced oversized, thus permitting greater detail and accuracy in the pattern.
6. True
7. ferric chloride, (FeCL_3).
8. False — The hardened resist protects the copper from being oxidized.
9. True — The circuit board maintains the sense and accuracy of the original circuit.
10. crossovers.
11. coupling.
12. False — Conductor widths vary over a range of sizes, and may be smaller than $\frac{1}{8}$ inch.
13. True
14. panel plating
15. False — Pattern plating requires less copper processing, and less copper is lost in the etching solution.
16. False — A tubular component should be cut from the board. The portion of the leads which remains should be used as terminal parts for the new component.
17. (b) metallic dust between the conductors.



1109A

1109A EXAMINATION CHECK SHEET

1. B - WHEN USING A PRINT-AND-ETCH METHOD TO PRODUCE A LARGE NUMBER OF BOARDS, ONE SHOULD USE -- A SILK SCREEN.

Of all the print-and-etch methods, the silk screen technique produces the best quality boards in the least amount of time.

2. D - WHEN USING A NEGATIVE-ACTING PHOTORESIST AND A PHOTO NEGATIVE CIRCUIT PATTERN, CONTACT PRINTING WILL CAUSE THE RESIST TO BECOME -- INSOLUBLE IN DEVELOPER UNDER THE LIGHT AREAS OF THE FILM NEGATIVE.

A negative-acting photoresist is initially soluble in a developer. Exposure to light through the clear areas of the film CHANGES the solubility such that this portion becomes insoluble in the developer.

3. B - A PLANNED JUNCTION OF TWO OR MORE CONDUCTORS MAY EXIST ON A PRINTED CIRCUIT BOARD ONLY -- IF AN ELECTRICAL CONNECTION IS INDICATED IN THE ORIGINAL SCHEMATIC. The crossing of conductors on a printed circuit board establishes an electrical connection. Conductors may be shown on a schematic diagram by crossing over other lines without indicating an electrical connection at that point.

4. A - EXCESSIVELY CLOSE SPACING BETWEEN CONDUCTORS MAY CAUSE A PROBLEM DURING COMPONENT ASSEMBLY WHICH IS KNOWN AS -- SOLDER BRIDGING. Solder bridging can occur even when unnoticed by the assembler, and can cause the circuit to fail.

5. A - ONE OF THE MAIN REASONS FOR USING PLATING PROCESSES IN CIRCUIT BOARD MANUFACTURING IS TO -- OBTAIN METALLIC BUILD-UP IN THE PRE-DRILLED HOLES OF THE BOARD. The process is generally known as THROUGH-HOLE PLATING.

6. B - THE RESIST MATERIAL USED IN THE FINAL ETCH CYCLE IN A PLATING PROCEDURE IS CALLED THE -- ETCH-RESIST METAL. The final etch is the cycle which removes all the unwanted copper. The solder plate which covers the circuit pattern and the through-holes is resistant to the etchant.

7. D - CIRCUIT TRACING PRINTED CIRCUIT BOARDS CAN BE MADE EASIER BY -- AVOIDING PARALLEL CIRCUIT CONFIGURATIONS. Provision should be made for connecting components, such as transistors in parallel, off the board through a connector or jumper wire so that each component can be tested separately.

8. B - A DEFECTIVE TUBULAR COMPONENT SHOULD BE REMOVED FROM A BOARD BY -- CUTTING THE BODY OF THE COMPONENT OUT. Cutting out the component ensures that minimum amount of damage will be done to the board itself.

9. D - A SMALL CRACK IN A PRINTED CIRCUIT CONDUCTOR SHOULD BE REPAIRED BY -- FLOWING A SMALL AMOUNT OF SOLDER OVER THE CRACK. When the crack is small and the conductor has not lifted, a small amount of solder or a short tinned (not stranded) wire can be used over the cracked area.

10. B - IF, IN TROUBLESHOOTING A PRINTED CIRCUIT BOARD, AN ELECTROLYTIC CAPACITOR IS SUSPECTED AS FAULTY BUT REPLACEMENT DOES NOT CORRECT THE PROBLEM, THE FAULT MAY BE DUE TO -- LEAKAGE RESISTANCE BETWEEN CONDUCTORS IN THE CAPACITOR CIRCUIT. Leakage resistance can be caused between conductors by metallic dust particles which may provide low-resistance discharge paths for capacitors.

PRACTICE EXERCISE SOLUTIONS

1. Printed circuits serve to support the components and provide the conductive paths between the components.
2. laminate.
3. resist.
4. silk screen.
5. False — The artwork is produced oversized, thus permitting greater detail and accuracy in the pattern.
6. True
7. ferric chloride, (FeCl_3).
8. False — The hardened resist protects the copper from being oxidized.
9. True — The circuit board maintains the sense and accuracy of the original circuit.
10. crossovers.
11. coupling.
12. False — Conductor widths vary over a range of sizes, and may be smaller than $\frac{1}{8}$ inch.
13. True
14. panel plating
15. False — Pattern plating requires less copper processing, and less copper is lost in the etching solution.
16. False — A tubular component should be cut from the board. The portion of the leads which remains should be used as terminal parts for the new component.
17. (b) metallic dust between the conductors.

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1109A

Example: Most window panes are made of
 A ☐ B ☐
 C ☒ D ☐
 A. wood. B. steel. C. glass. D. iron.

1. A ☐ B ☐ C ☐ D ☐ **When using a print-and-etch method to produce a large number of boards, one should use**
 (A) a negative-acting photoresist. (B) a silk screen. (C) tape-resist applied directly to the cladding. (D) liquid resist painted directly on the cladding.

2. A ☐ B ☐ C ☐ D ☐ **When using a negative-acting photoresist and a photo negative circuit pattern, contact printing will cause the resist to become**
 (A) soluble in developer under the light areas of the film negative. (B) insoluble in developer under the dark areas of the film negative. (C) soluble in developer under the dark areas of the film negative. (D) insoluble in developer under the light areas of the film negative.

3. A ☐ B ☐ C ☐ D ☐ **A planned junction of two or more conductors may exist on a printed circuit board only**
 (A) if jumper wires cannot be used. (B) if an electrical connection is indicated in the original schematic. (C) on double-sided circuit boards. (D) in order to simplify the layout.

4. A ☐ B ☐ C ☐ D ☐ **Excessively close spacing between conductors may cause a problem during component assembly which is known as**
 (A) solder bridging. (B) pinholing. (C) heat sinking. (D) power loss.

5. A ☐ B ☐ C ☐ D ☐ **One of the main reasons for using plating processes in circuit board manufacturing is to**
 (A) obtain metallic build-up in the pre-drilled holes of the board. (B) avoid conductor crossovers. (C) reduce the number of processing steps in producing a printed circuit board. (D) to lower the cost of producing a printed circuit board.

6. A ☐ B ☐ C ☐ D ☐ **The resist material used in the final etch cycle in a plating procedure is called the**
 (A) photoresist. (B) etch-resist metal. (C) vapor plating. (D) stripping.

7. A ☐ B ☐ C ☐ D ☐ **Circuit tracing printed circuit boards can be made easier by**
 (A) avoiding the use of crossovers. (B) using curved conductor shapes rather than straight, angled conductor shapes. (C) heavily loading the output circuit to see if the board fails. (D) avoiding parallel component configurations.

8. A ☐ B ☐ C ☐ D ☐ **A defective tubular component should be removed from a board by**
 (A) heating the component leads until the part comes out easily. (B) cutting the body of the component out. (C) cutting the twisted or crimped leads on the underside of the board. (D) straightening the lugs on the underside of the board and lifting the part out.

9. A ☐ B ☐ C ☐ D ☐ **A small crack in a printed circuit conductor should be repaired by**
 (A) severing the conductor with a sharp knife and removing the faulty section. (B) soldering a small length of stranded wire across the crack. (C) soldering a jumper wire from the component to the nearest available terminal. (D) flowing a small amount of solder over the crack.

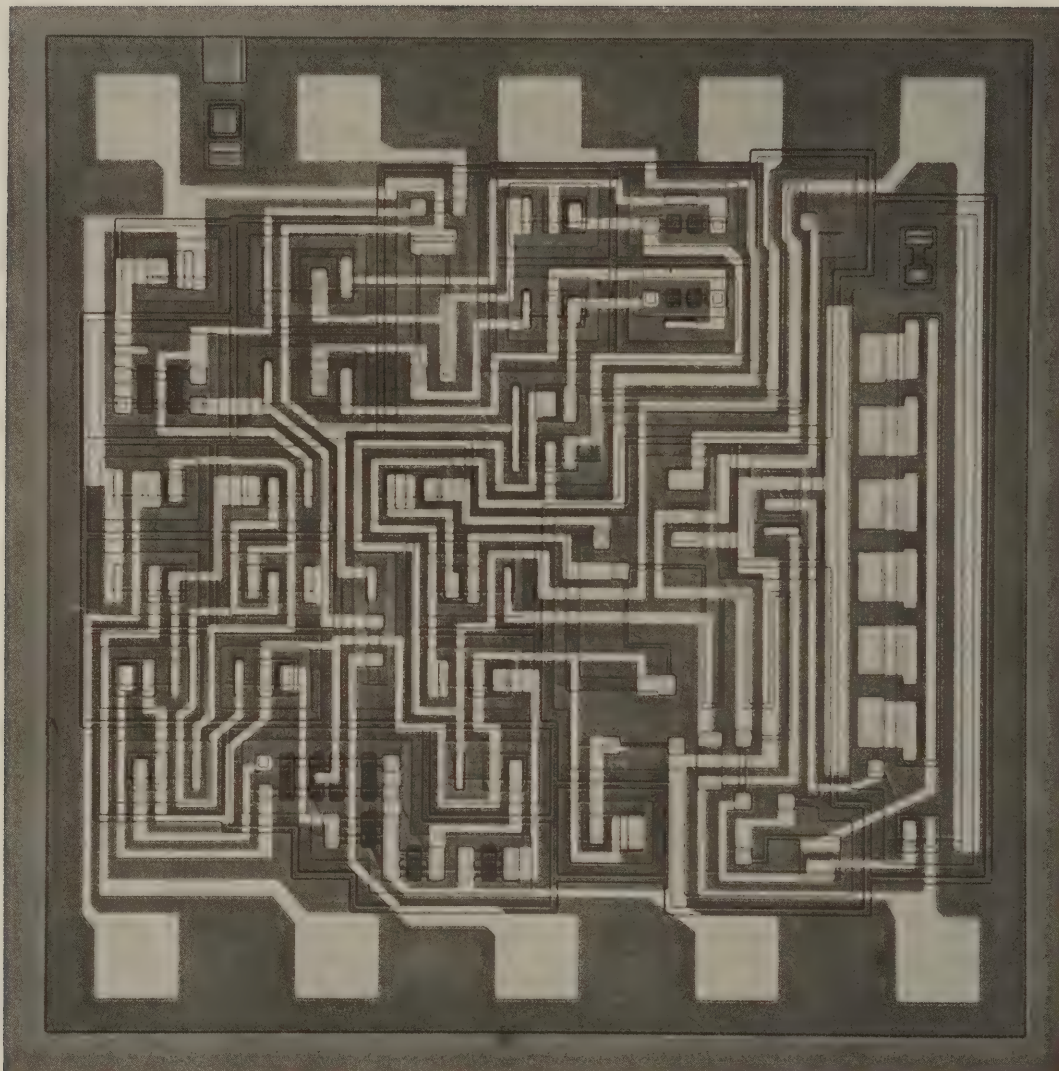
10. A ☐ B ☐ C ☐ D ☐ **If, in troubleshooting a printed circuit board, an electrolytic capacitor is suspected as shorted but replacement does not correct the problem, the fault may be due to**
 (A) a partially-cracked printed circuit conductor. (B) leakage resistance between conductors in the capacitor circuit. (C) a high-resistance connection (greater than 1,000 megohms). (D) a severed (open) printed circuit conductor.

INTEGRATED CIRCUITS

1112

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





Monolithic integrated circuits, when viewed through a microscope, appear as shown in the photo micrograph above. The light areas are metallized layers which interconnect the various parts of the integrated circuit.

Courtesy Fairchild Semiconductor

CONTENTS

Basic Construction	Page 6
Monolithic Integrated Circuits	Page 8
Crystal Growth	Page 9
Circuit Fabrication	Page 13
Hybrid Integrated Circuits	Page 21
Component Limitations	Page 22
Linear Circuits	Page 23
Differential Amplifiers	Page 24
Operational Amplifiers	Page 25
Digital Circuits	Page 27
Resistor-Transistor Logic (RTL)	Page 28
Diode-Transistor Logic (DTL)	Page 30
Transistor-Transistor Logic (TTL)	Page 33

*The happy people are those who are producing something; the
bored people are those who are consuming much and producing
nothing.*

—Dean W. R. Inge

INTEGRATED CIRCUITS

Integrated circuits are a part of the broad field of MICROELECTRONICS. There are several categories under the general heading of microelectronics as shown by the chart of Figure 1. The main divisions are **DISCRETE COMPONENTS** and **INTEGRATED CIRCUITS**.

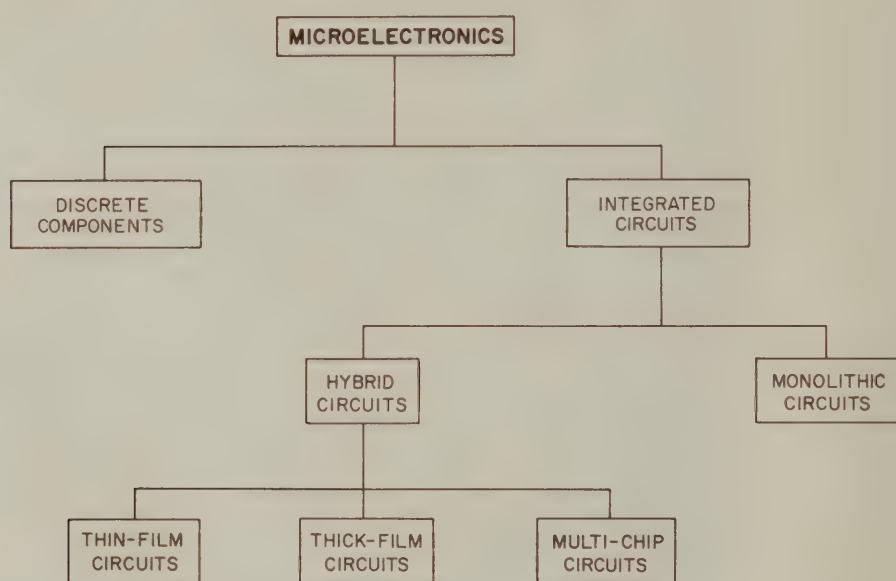


Figure 1

~~Discrete components are individually packaged devices such as transistors and diodes. They are called discrete components because they are single devices within themselves. There are no additional components such as resistors, diodes, or other transistors within the transistor case. Integrated circuits, however, do contain more than one element or device within the same package. An integrated circuit can be designed to perform a complete circuit function. Integrated circuits can be constructed as monolithic, hybrid, thick film or thin film circuits. These techniques will be covered later in this lesson.~~

Discrete solid state semiconductor components and integrated circuits have much in common. Both are made of crystalline materials such as silicon and germanium. Thus, there are similarities in the construction techniques between both types. Before examining integrated circuits, let's review the basics of solid state electronics.

Transistors and integrated circuits are both made of SEMICONDUCTOR materials. A SEMICONDUCTOR MATERIAL is one which is neither a good conductor of electricity such as copper, nor a nonconductor such as mica. Instead, the material lies somewhere between the metallic conductors and the insulators in conductivity. The CONDUCTIVITY of a material is determined by the number of FREE OR EXCESS ELECTRONS present in the material.

The structure of pure silicon and germanium is CRYSTALLINE, and the molecules are in an ordered array such as shown in Figure 2. This arrangement is called a DIAMOND LATTICE since the lattice-like structure is similar to that found in high-quality diamond crystals. Note that there are no free or excess electrons in the silicon lattice structure of Figure 2. Each atom is balanced with just four (the required number) electrons. With all the electrons held by the lattice structure, no excess electrons are free to drift through the crystal. Theoretically, this represents a completely stable lattice structure and would be a perfect insulator. The crystal would not be capable of conducting an electrical current.

However, such perfect crystals do not exist in actual practice. Even in highly pure crystals some free electrons are present due to impurities in the material, thus making the crystal a poor conductor rather than a nonconductor. These imperfections are often due to small IMPURITIES such as trace amounts of dirt, metal, or other foreign substances. IMPURITIES PLAY AN EXTREMELY IMPORTANT ROLE IN SEMICONDUCTOR THEORY.

Since impurities can cause an insulator to become a partial conductor, the PURPOSEFUL ADDITION of some kinds of impurities can cause the crystal to become a relatively good conductor. However, the type of impurity is not the only consideration. A very important factor in the manufacture of a transistor is the CONCENTRATION of the added impurity. When impurities are added to the silicon or germanium, the process must be carefully controlled. The process of adding impurities to the silicon or germanium used in a transistor is called DOPING.

One type of doping procedure adds free electrons to a crystal. The impurity which enters the lattice structure is a substance in which the atoms have one more electron than silicon atoms have. PHOSPHOROUS, ARSENIC and ANTIMONY are typical impurities which will add free electrons to a silicon lattice. For this reason, phosphorous, arsenic and antimony are called DONOR ELEMENTS, and they form N-TYPE semiconductors (negative charge). Figure 3A shows a lattice structure which contains phosphorous impurity atoms.

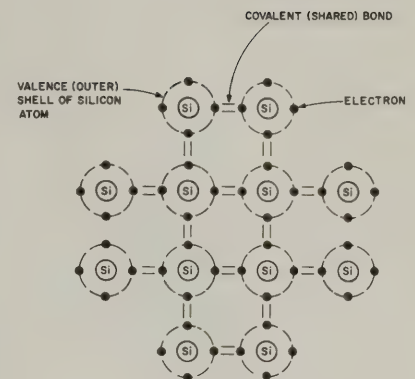


Figure
2

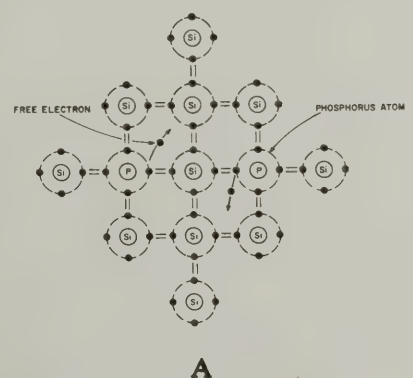


Figure
3

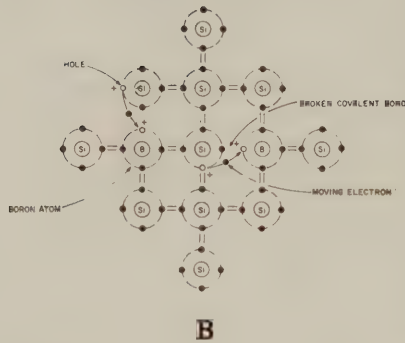


Figure 3

Another type of doping procedure adds atoms to the crystal structure which are DEFICIENT in electrons. The atoms of BORON, ALUMINUM, GALLIUM and INDIUM are impurities which are deficient in electrons. Since these atoms have one fewer electron than the "perfect" structure, they can ACCEPT a free electron. Thus, boron, aluminum, gallium and indium are called ACCEPTOR ELEMENTS, and form P-TYPE semiconductors (positive charge). It is possible to think of these deficient atoms as having "holes" where a crystal-lattice electron should be. A lattice structure containing boron impurity atoms is shown in Figure 3B.

The concept of free holes and free electrons becomes clearer upon returning briefly to the discussion of conduction. As mentioned previously, highly purified silicon or germanium is a very poor conductor. Transistors, however, require crystalline material of greater conductivity than that of pure silicon or germanium. The conductivity of crystalline material can be increased by either of two methods: (1) by adding impurities to the crystal, as shown earlier, or (2) by HEATING the crystal. When a silicon crystal is heated sufficiently, the atoms acquire additional energy, causing them to vibrate more violently. This vibration tends to disrupt the crystal lattice structure by breaking some of the electrons away from their ordered positions. When this happens, the electrons become mobile and move about in the crystal in random fashion as shown in Figure 4. Thus, a particular electron moves from atom to atom, constantly taking the place of a different electron.

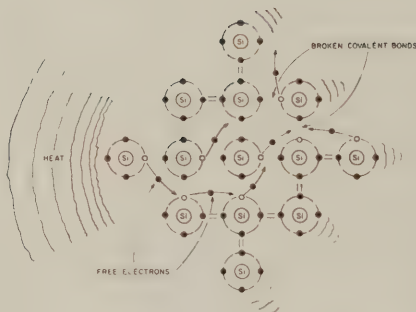


Figure 4

A similar situation exists in the case of free holes. Note that each time an electron breaks away from a silicon atom, it leaves a "hole". This process is called the THERMAL GENERATION OF HOLE-ELECTRON PAIRS. The process of continually moving electrons produces the EFFECT of continuously moving holes. When an electron and a hole meet, or come together, the electron fills the electron deficiency of the atom. At this time, both the electron and the hole cease to exist as free charge carriers. This process is called RECOMBINATION.

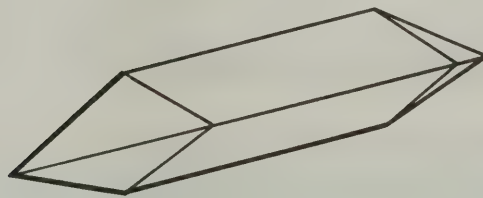
Electrons and holes are CHARGE CARRIERS in electric circuit conduction. Electrons are negatively charged, and holes are positively charged. The process of recombination in a diode or transistor is controlled, however, and not random as in thermal generation. Conduction in a transistor is caused by VOLTAGE DIFFERENCES across N-type and P-type semiconductors.

BASIC CONSTRUCTION

Integrated circuit fabrication techniques have developed out of transistor technology. The process which is most basic to both is the DIFFUSION DOPING procedure which introduces impurities, and thus produces P or N-type silicon or germanium.

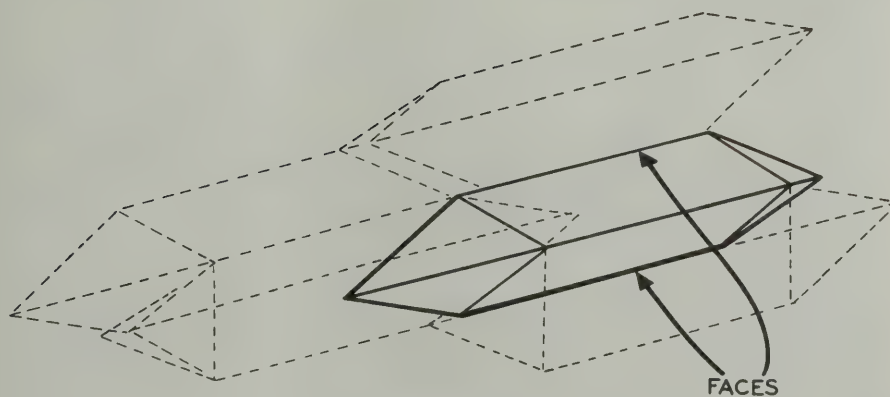
Integrated circuits contain a number of electrical components constructed on a common base, called a **SUBSTRATE**. This base is a slice or wafer of single crystal silicon or germanium. It is important to note that the various P and N regions of a transistor or integrated circuit are **MONOCRYSTALLINE**. That is, the silicon or germanium is "grown" as a single crystal.

Quartz, germanium and a variety of other minerals have crystalline form. Each mineral possesses its own unique crystal shape, and always forms this shape under natural growing conditions. Figure 5A shows a typical crystal. Note that it is basically prism-shaped, and also slightly irregular. Crystals such as this, occurring naturally, are very small **MONOCRYSTALS**. This means that the lattice structure forms in the same continuous direction throughout the crystal. Crystals possess the capability of maintaining their characteristic lattice structure during formation in a continuous "perfect" lattice. As the crystal continues to form, however, the lattice may begin to



MONOCRYSTAL

A



POLYCRYSTAL

B

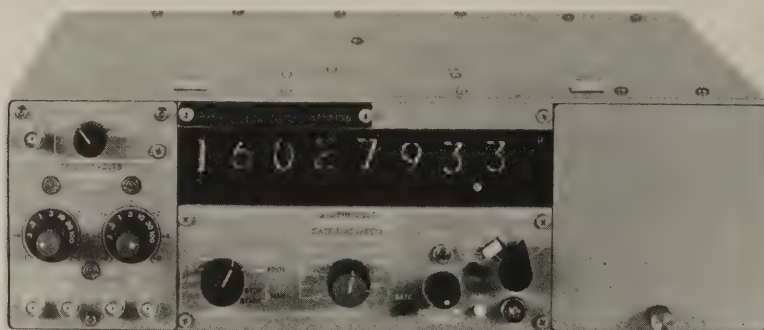
Figure 5

form in another DIRECTION while maintaining its characteristic structure. Boundaries, or FACES, are then formed between the lattice structures where this change in direction occurs. This boundary marks the division between many side-by-side SEPARATE CRYSTALS, as shown in Figure 5B. All the crystals are physically identical, but they are oriented differently with respect to each other. Thus, what appears to be a single crystal of silicon, germanium or quartz is actually many single crystals formed together. These minerals are said to possess POLYCRYSTALLINE form.

Polycrystalline material is not suitable for use in semiconductor electronics because the boundaries between crystals interfere with the conduction of charge carriers. For this reason, semiconductor material is grown as a single, large monocrystal, and later cut into thin slices to form the wafer or substrate. Since this material is monocrystalline, the lattice structure is uniform and conduction characteristics can be accurately predicted. Several methods are used in integrated circuit technology for producing monocrystalline wafers, slices and layers. These will be discussed as the various integrated circuit construction processes are explained.

Monolithic Integrated Circuits

A MONOLITHIC INTEGRATED CIRCUIT is constructed on a single piece of silicon material in much the same manner as a single transistor, and is the simplest of all integrated circuits. The main difference between them is the number of devices which are made on the same base, or substrate. The process by which the P and N-type regions of a transistor are formed are the same as those for producing passive components (resistors and capacitors) on the substrate. Thus, transistors, resistors and capacitors may be formed on a single substrate simultaneously. Once formed, the components on a



Digital integrated circuits are widely used in electronic counters like the one shown above. Such an instrument, using discrete components, would be many times the size of the IC unit.

Courtesy Computer Measurements Corp.

monolithic circuit cannot be removed separately for repair. When any single component fails, the entire integrated circuit must be replaced.

Crystal Growth

The first step in producing an integrated circuit is to produce the substrate on which the transistor is formed. This is much the same as obtaining a CHASSIS on which to place the components in a conventional piece of electronic equipment. Producing the substrate begins with a semiconductor MELT of silicon. A semiconductor melt is a container of molten silicon to which has been added a predetermined amount of P or N-type impurities.

A very small single crystal, called a SEED, is used to "grow" a large CRYSTAL INGOT. This is shown in Figure 6. The large crystal is grown in such a way as to maintain the unidirectional crystal lattice structure of the SEED crystal. The seed crystal is placed in the semiconductor melt and slowly withdrawn while being rotated. By controlling the rotation and temperature, the molten material adheres to the seed, and the crystal grows in size. A crystal ingot of one or two inches in diameter and several inches in length is produced in this way. More important, however, is the fact that

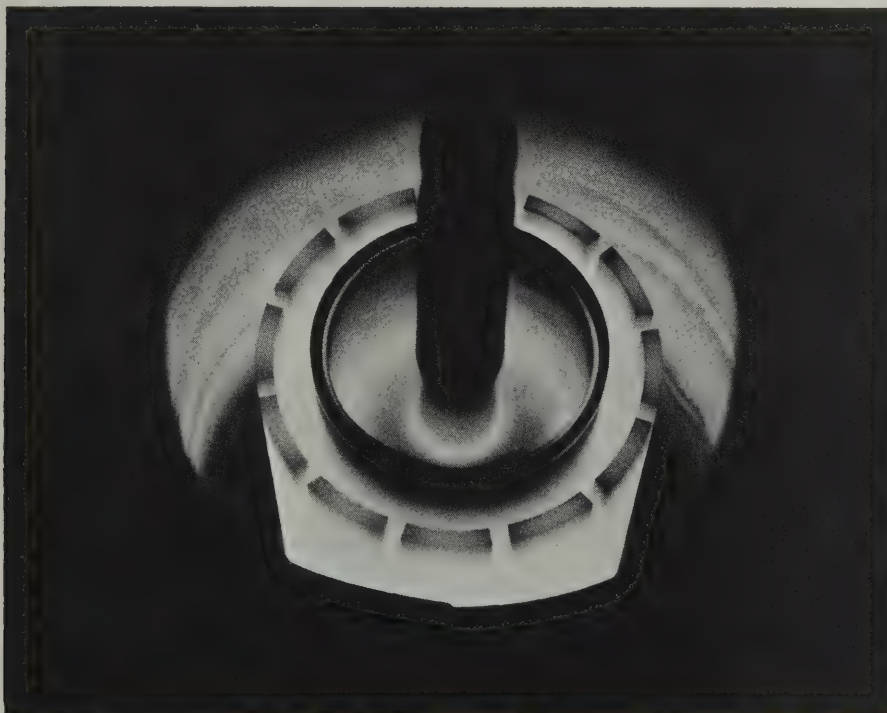


Figure 6

the ingot is a **SINGLE CRYSTAL**. The crystal lattice is oriented in the same direction throughout the ingot with no boundaries or faces as with a polycrystal.

Thus we see that one method for producing P or N-doped semiconductor material is to produce it **AS THE CRYSTAL IS FORMED** from a semiconductor melt.

The crystal ingot is then sliced into thin wafers, each of which becomes a substrate used in making integrated circuits. The wafers have, then, the same diameter as the crystal ingot—one to two inches. Although there are several subsequent steps which may vary depending upon the type of integrated circuit, the crystal growth substrate procedures are quite uniform. The substrate produced from an impurity melt is most often P-type material. That is, P-type impurities are added to the molten silicon, and thus the wafers are P-type. One side of the wafer is then polished to a mirror-smooth finish in preparation for subsequent processing. The final wafer thickness is usually less than 10 mils thick.

The next step in preparing the substrate is to add a very thin layer of N-type material to the smooth side of the P-type wafer. This is done by another

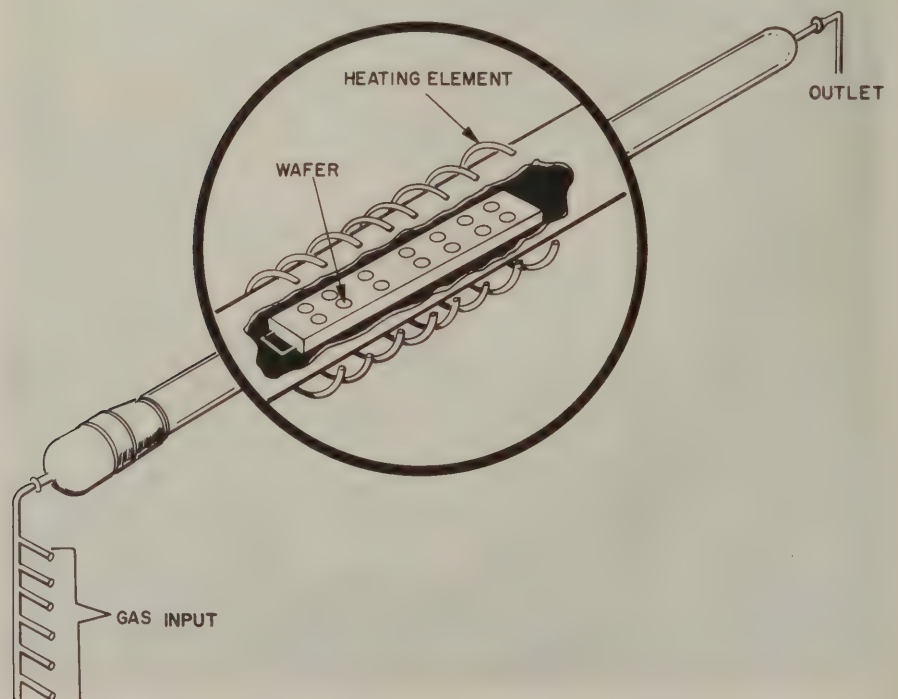


Figure 7

method for producing impurity doped silicon semiconductor material called **EPITAXIAL GROWTH**. Like the impurity melt method, the crystal material is grown with the proper impurities introduced as the crystal is growing. An epitaxially-grown layer is not produced from a melt, however. This type of crystal is grown from the gaseous state directly to the solid crystal state. A number of wafers are placed in an EPITAXIAL REACTOR which is essentially a long glass tube, as shown in Figure 7.

The tube is surrounded with a heating element and the temperature is precisely controlled. The proper mixture of gases is admitted at one end of the tube and vented at the other end through an outlet port. As the gas passes over the wafers, a doped silicon layer forms on the wafer surfaces. This epitaxial layer is doped with impurities of the opposite type from the wafer on which it is grown. When the proper layer thickness has been reached, the silicon gases are stopped. The wafer is then as shown in Figure 8. The boundary between the original P-type wafer and the N-type epitaxial layer forms a P-N JUNCTION.

Certain aspects of integrated circuit fabrication are similar to printed circuit fabrication. The circuit pattern must be placed onto the substrate in a very

EPITAXIAL WAFER

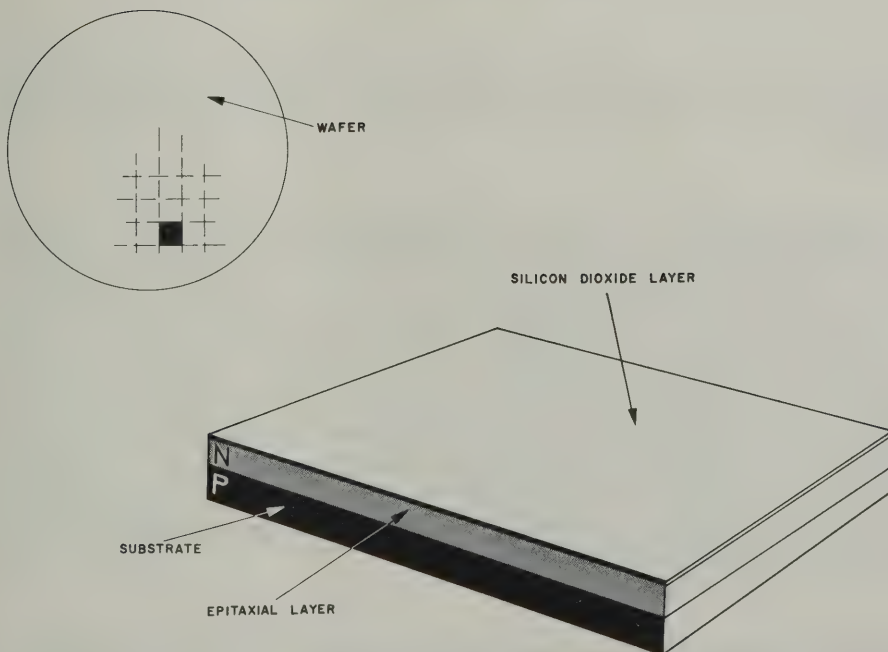


Figure 8

An epitaxial wafer, the starting point for an epitaxial-diffused integrated circuit.

precise manner. Some very complex equipment is required to do this. However, it is the process itself which is important, and we will neglect the detailed operation of the equipment.

The circuit pattern construction, in the case of integrated circuits, consists of several separate elements. One of these is the production of transistor collector, emitter and base areas. Another is the production of passive components (resistors and capacitors). The components, once produced, are connected together into a working circuit.

An important part of integrated circuit fabrication is the process of **DIFFUSION**. Transistors are made of P and N-type silicon or germanium. P and N-type silicon was shown to be produced either as the crystal is grown from a MELT, or by growing an EPITAXIAL LAYER. Diffusion is another method for producing P and N-type silicon by introducing impurity atoms of either P or N material into a silicon wafer.

One method for diffusing impurity atoms into a silicon wafer is to heat the wafer to approximately 1000°C in a furnace. The heated wafer is then subjected to a gas flow containing a heavy concentration of impurity atoms. The heating tends to make the silicon wafer temporarily more porous, and allows some of the impurity atoms to penetrate the silicon. The impurity atoms then become part of the lattice structure.

By controlling the temperature, the diffusion time and the concentration of impurity gas, the impurity atoms penetrate the silicon wafer to the desired depth.

Diffusion doping has an additional capability over the other types of doping. Assume that a silicon wafer has been doped with N-type impurities by one of the doping methods discussed so far. If P-type impurities are diffused into the wafer, the N-type impurities will be "neutralized". When all of the original doping has been neutralized, the wafer behaves as pure, or intrinsic, silicon. By continuing the diffusion, the original N-type wafer can be converted to a P-type wafer. This is accomplished by continuing the P-type diffusion until the P-type concentration becomes greater than the original N-type concentration. Thus, through the diffusion process it is possible to convert intrinsic silicon to either P or N-type, or to convert one type of doped silicon to the other type.

The impurity gas stream containing the diffusant causes a heavy impurity buildup near the surface of the wafer called the SURFACE CONCENTRATION. These impurity atoms spread deeper into the wafer as the diffusion

progresses. A P-N junction is formed at the depth where the diffused impurity concentration just equals the concentration of the opposite impurity atoms already in the wafer. A P-N junction produced by the diffusion process is physically similar to a P-N junction produced by the epitaxial process. Both methods produce a wafer similar to the one shown in Figure 8.

The diffusion process is most useful when used to produce selectively doped areas on a single wafer. These doped areas are produced in a series of separate diffusions.

Circuit Fabrication

The first step in producing the actual integrated circuit is to cover the silicon wafer with a protective coating which is immune to impurity penetrations at high temperatures. Fortunately, silicon dioxide (SiO_2) is a material which is both immune to impurity penetrations and is also easily produced by simply placing the wafer in an oxygen atmosphere. This covers the entire wafer with a thin layer of SiO_2 . However, in order to permit diffusions only in certain

ISOLATION MASKING

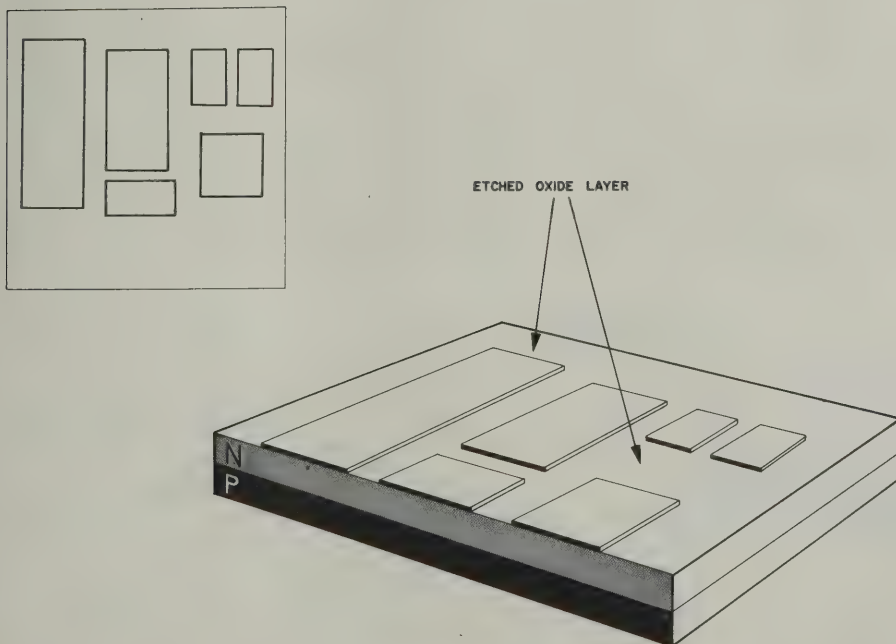


Figure 9

The wafer with the isolation masking pattern etched into the SiO_2 layer.

areas, it is necessary to etch away the SiO_2 layer in the regions where diffusion is to take place, as shown in Figure 9. The areas which are not removed by etching prevent diffusion into the silicon below. The resulting SELECTIVE DIFFUSIONS are used to form the collector, base and emitter areas of transistors, and also the passive components which are located on the same wafer.

Since the SiO_2 layer is used as a mask for the silicon, some method is needed for selectively removing the portions of the SiO_2 layer. In order to do this, another mask, called a PHOTOMASK, is used which permits SELECTIVE ETCHING of the SiO_2 layer. This mask is used in a PHOTOETCH method similar to that used in printed circuit board production.

The photomask is produced from a greatly enlarged ARTWORK MASTER. This master is a pattern of the areas on the substrate which will receive that diffusion. Since the same wafer receives more than one diffusion, each piece of artwork must be highly accurate. The pattern is therefore generally produced on a precision mechanical drafting table, called a COORDINATOGRAPH. The large pattern is photographically reduced in size. The network pattern is thus faithfully reproduced, but appears as a microscopically small section on a piece of photographic film.

As pointed out, the silicon wafer, which is used as a substrate for the integrated circuit, is one to two inches in diameter. This is much larger than required for just one integrated circuit. Therefore, the final mask contains not only one small circuit pattern, but perhaps hundreds of identical patterns placed very close together. In another precision device, called a step-and-repeat machine, the photo-reduced pattern is contact printed over and over throughout the surface area on a piece of film slightly larger than the wafer itself. The multiple images are transferred to a photo-sensitive glass plate. The glass plate is the FINAL MASK, and is used to produce a photoprint

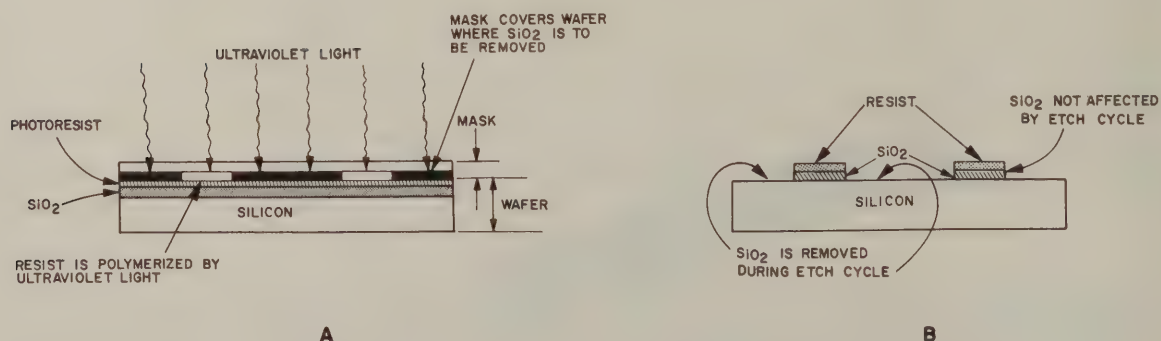


Figure 10

of the patterns on the wafer. The figures illustrating the following steps show only one of the circuits on the wafer. It must be remembered, however, that the fabrication steps are performed on all of the circuits simultaneously.

The photoprint and SiO_2 etch procedure is also similar to the method used in printed circuit manufacture and is shown in Figure 10. The wafer is coated with a ~~PHOTOSENSITIVE RESIST~~ (photoresist). A photoresist is a polymer resin material which is sensitive to ultraviolet (UV) light. Exposure to UV changes the solubility characteristics of the photoresist. This means that a thin coating of photoresist is applied to the wafer which hardens, or polymerizes, upon exposure to UV. By laying the glass mask over the wafer, only the portions of the photoresist under the clear areas of the mask are exposed. This changes the solubility of these areas and makes them insoluble in a developer solvent. The areas under the dark portions of the mask remain soluble during the exposure period. Thus, the developer REMOVES the resist in the unexposed areas and completely uncovers the SiO_2 beneath.

The next step is selective etching—removal of the SiO_2 where the resist was removed. The etching solution attacks and removes the uncovered SiO_2 , but

FIRST DIFFUSION

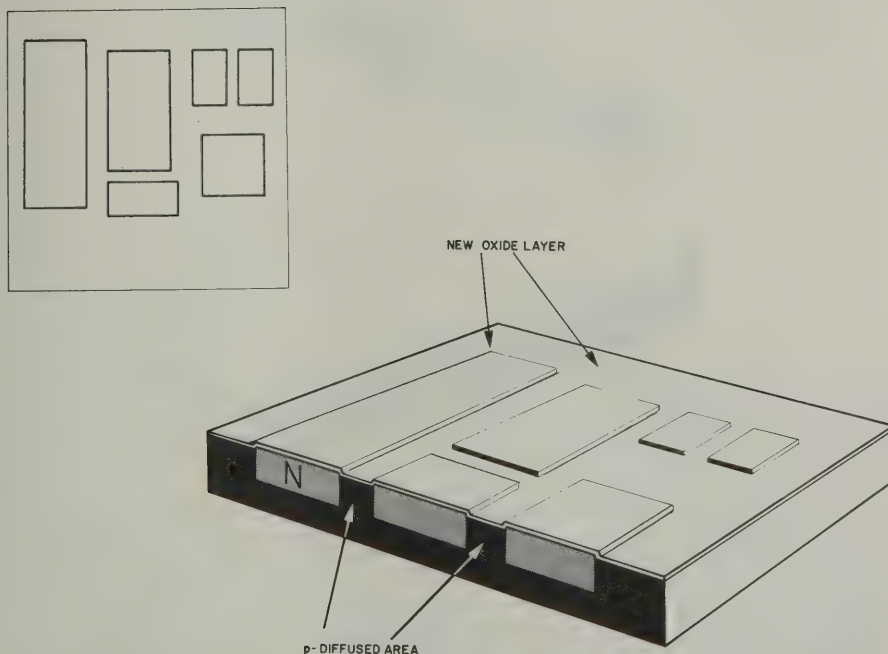


Figure 11
The wafer after the first (a P-type) diffusion.

does not affect the hardened resist. Since the hardened resist is not removed, the SiO_2 beneath it is not removed in these areas, as shown in Figure 10B.

When the etching is completed, the first diffusion may be performed. This is a deep P-type diffusion which is allowed to completely permeate the N-type epitaxial layer and meet the P-type substrate below. The P-type impurities enter the silicon, however, only where the SiO_2 layer has been etched away. The result of the first diffusion is that a number of N-type "islands" are left in a solid P-type silicon block. When the diffusion is complete, the hardened resist, which was unaffected by the developer and the etching solution is completely removed. This is often done by scrubbing the surface with heated solvents. Finally, a layer of SiO_2 is REGROWN over the entire wafer. This is done in preparation for another diffusion step similar to the first. For this diffusion step it is necessary to begin with a complete, unbroken layer of SiO_2 over the wafer. At this point the wafer appears as shown in Figure 11.

RESISTOR/BASE/ANODE MASKING

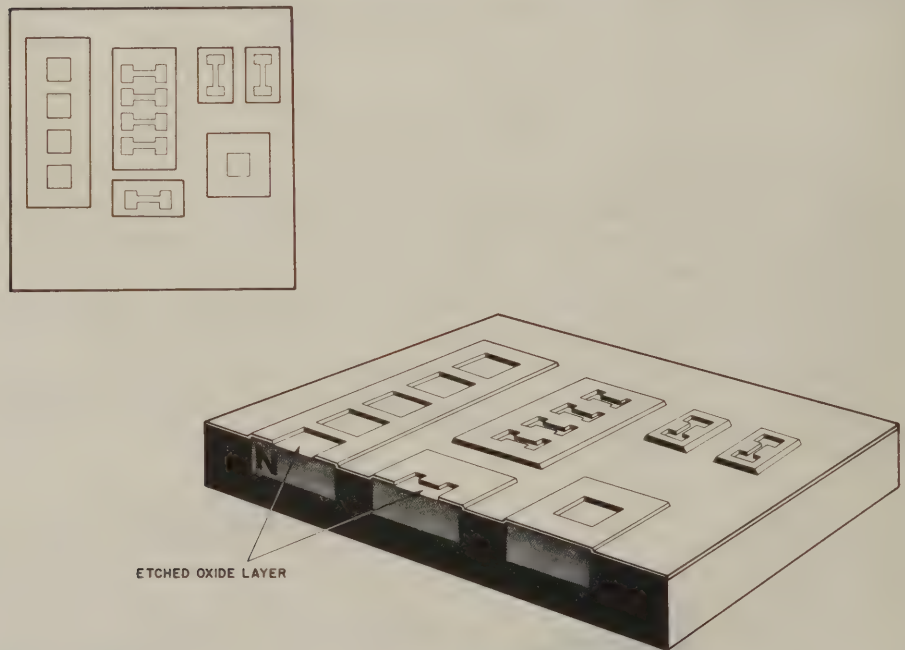


Figure 12
The wafer with the second pattern etched into the oxide layer.

The second diffusion requires a new SiO_2 mask. This mask is prepared as before from a new artwork pattern and transferred to a glass plate. The wafer is again coated with a photoresist, exposed to ultraviolet light, and

developed. The development cycle removes the resist coating only in the areas which are to receive diffusions. The wafer now appears as shown in Figure 12. This time the diffusions are quite shallow, and are made only in the N-type islands. The second diffusion is again a P-type diffusion, and it forms transistor base regions, resistors and the anodes of diodes and capacitors. When the second diffusion is complete, the remaining resist is removed and another layer of SiO_2 is grown over the complete wafer, as shown in Figure 13.

SECOND DIFFUSION

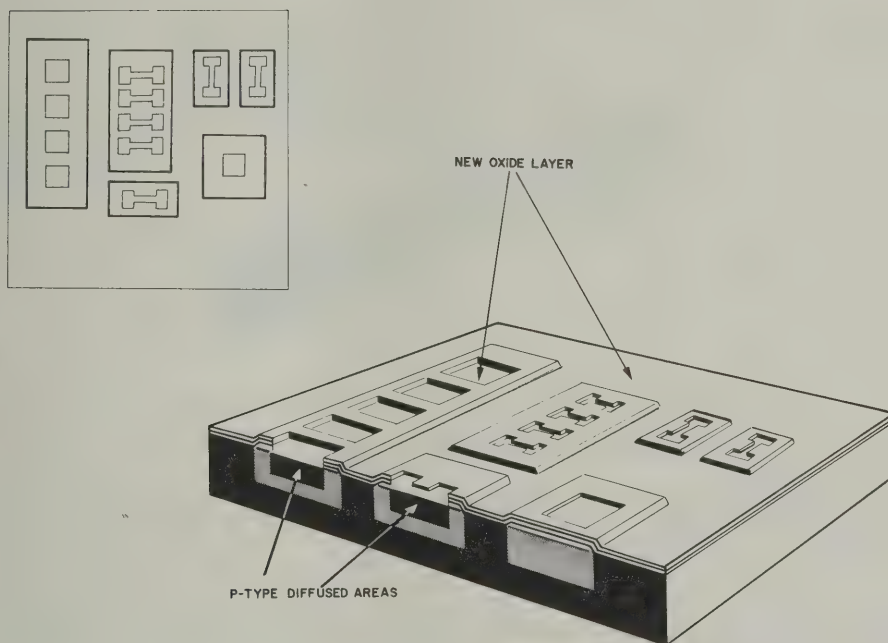


Figure 13
The wafer after the second (again P-type) diffusion.

The third diffusion, using a third mask, is accomplished just as the others, but is located in the P-diffused base regions. The wafer now appears as shown in Figure 14. The third diffusion is N-TYPE. This forms the transistor emitters and cathode regions for diodes and capacitors. A narrow groove is also diffused into the original N-type islands. This is done to provide a contacting area to the N-type collector areas. After the third diffusion is complete, the wafer appears as shown in Figure 15. Again, an SiO_2 layer is grown over the wafer. The formation of junction areas is complete after the third diffusion. Photoresist is again applied to the wafer surface. A glass mask is then placed over the photoresist and ultraviolet light is used to expose the photoresist. Windows are etched through the SiO_2 where connections are to be made. A mask is required to produce the METALLIZATION, or

EMITTER/CATHODE MASKING

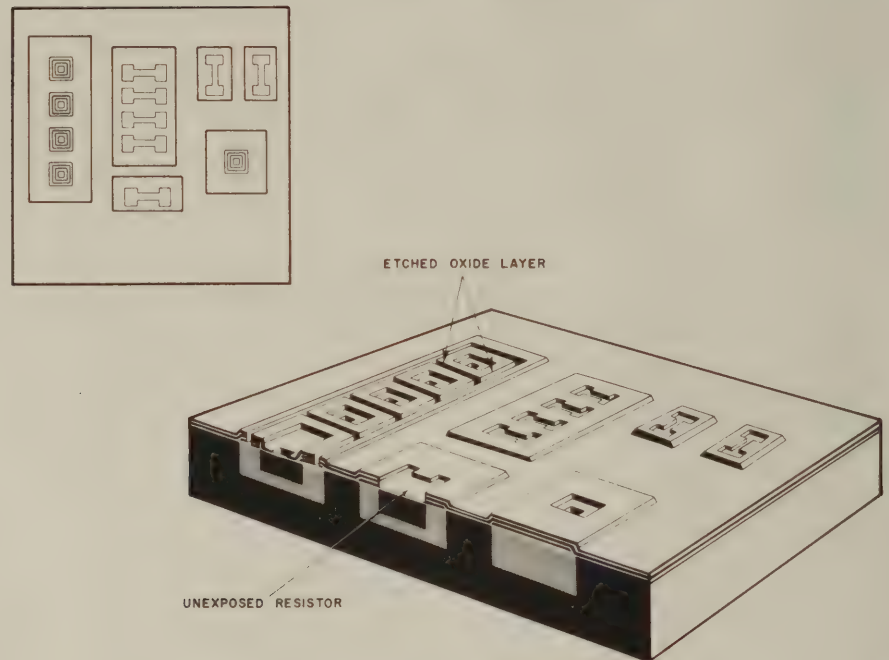


Figure 14
The wafer with the emitter/cathode pattern etched into the SiO_2 layer.

metallic aluminum interconnections between the components produced by diffusions. This step is shown in Figure 16.

To produce the metallized contacts, a thin coating of aluminum is vacuum-deposited over the entire surface of the wafer. The unwanted aluminum is etched away using this mask and photoresist techniques. The remaining metal provides the required interconnection between the integrated circuit components.

When the metallization step is completed, the wafer appears as in Figure 17, and production of the actual integrated circuit is completed. It must be remembered, however, that each mask contains MULTIPLE pattern images, and many circuits are produced simultaneously on a single wafer. The multiple-circuit wafer is shown in Figure 18. The remaining task, then, is to separate the individual circuits. This is done by dragging a diamond-tipped stylus in a crosshatch pattern between the circuits. This procedure is called SCRIBING. The scribed lines are actually scratches in the silicon. The silicon breaks along these lines when properly stressed. The result is that many individual and identical circuit DICE are produced. A single circuit, or DIE, is mounted and packaged in one of several available ways, ready to be used in a working circuit, as shown in Figure 19.

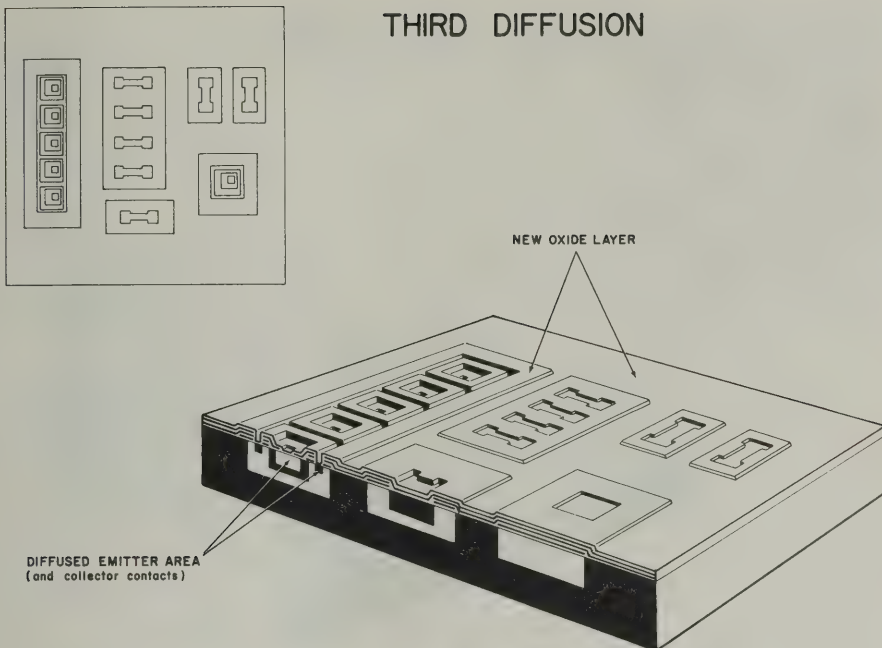


Figure 15
The wafer after the third (N-type) diffusion.

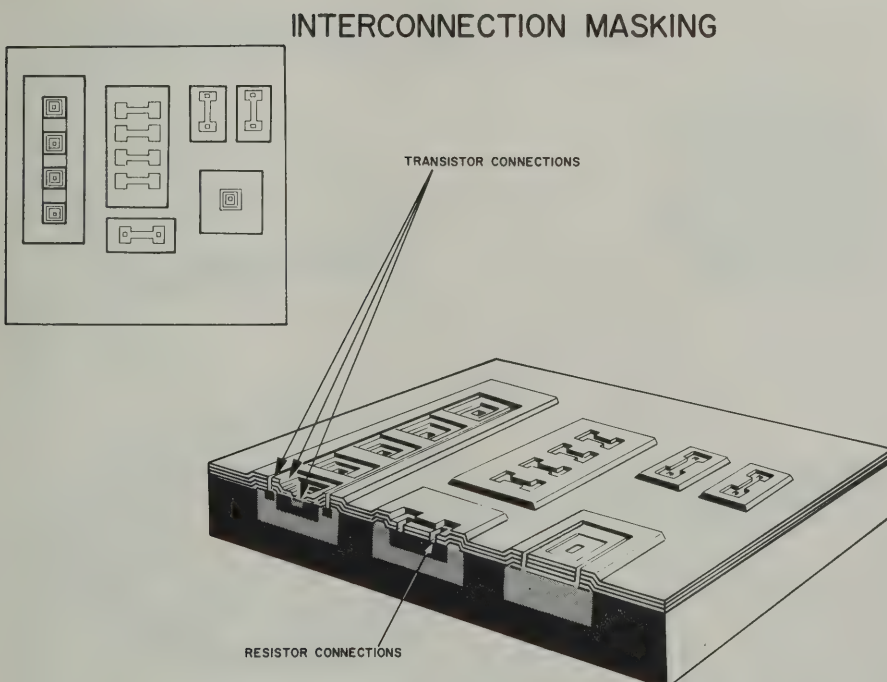
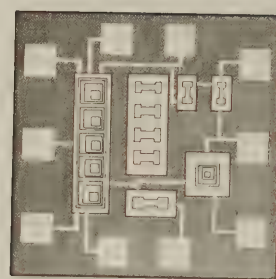


Figure 16
The wafer with openings etched into the SiO_2 layer to permit electrical contact to the various components of the integrated circuit.



METALLIZATION

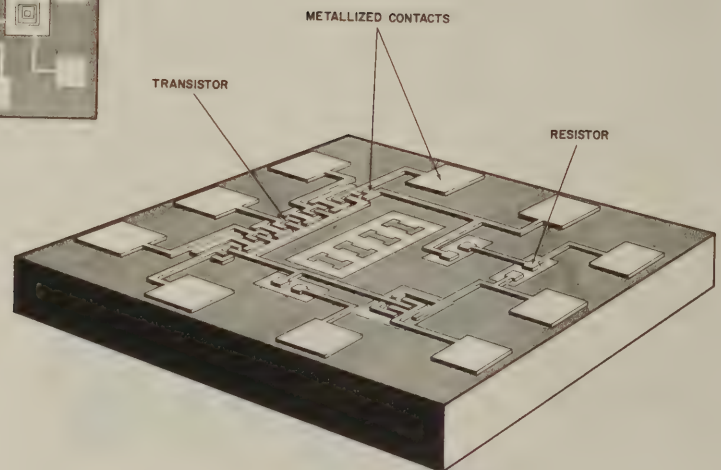


Figure 17
The wafer with the metallized contacts deposited.

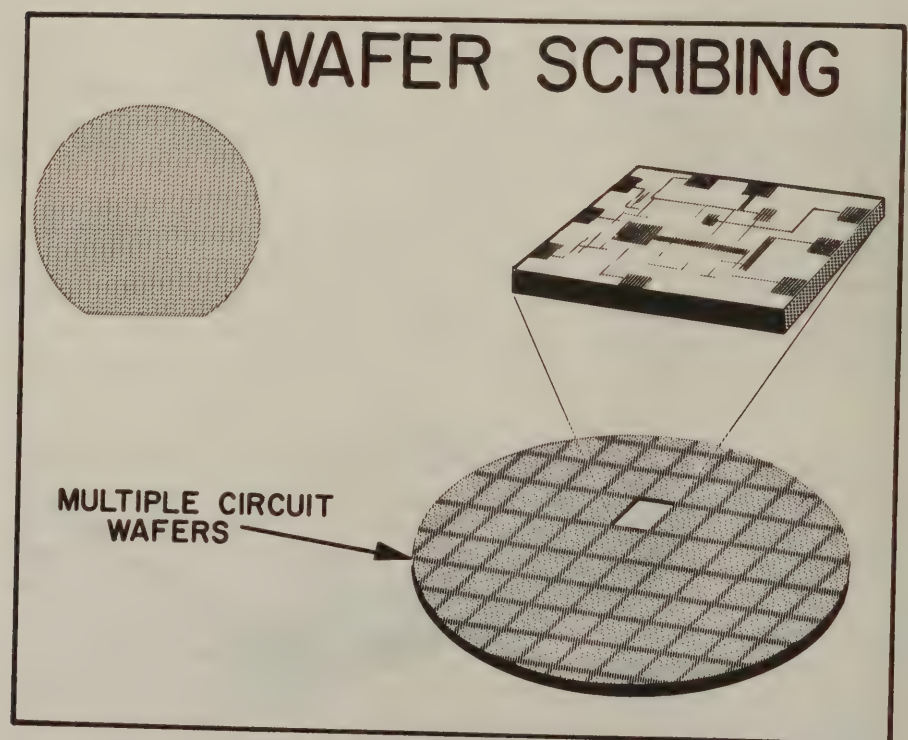


Figure 18
Separation of the wafer into the individual circuits by scribing.

DIE AND WIRE BONDING

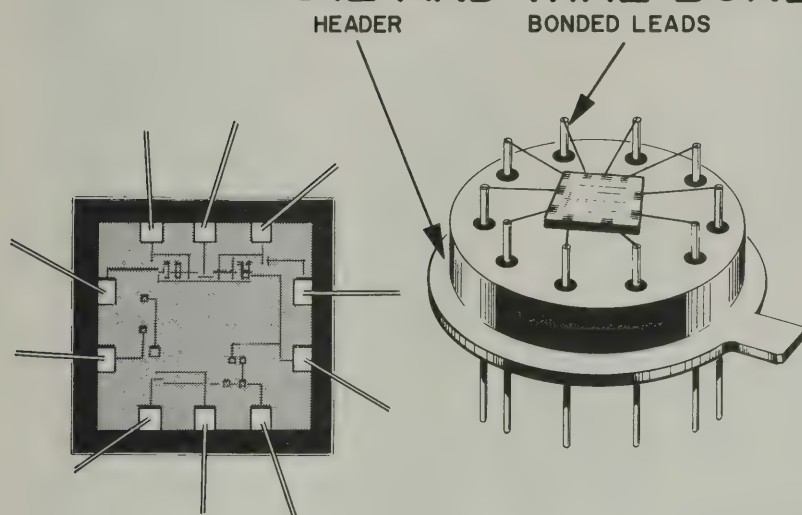


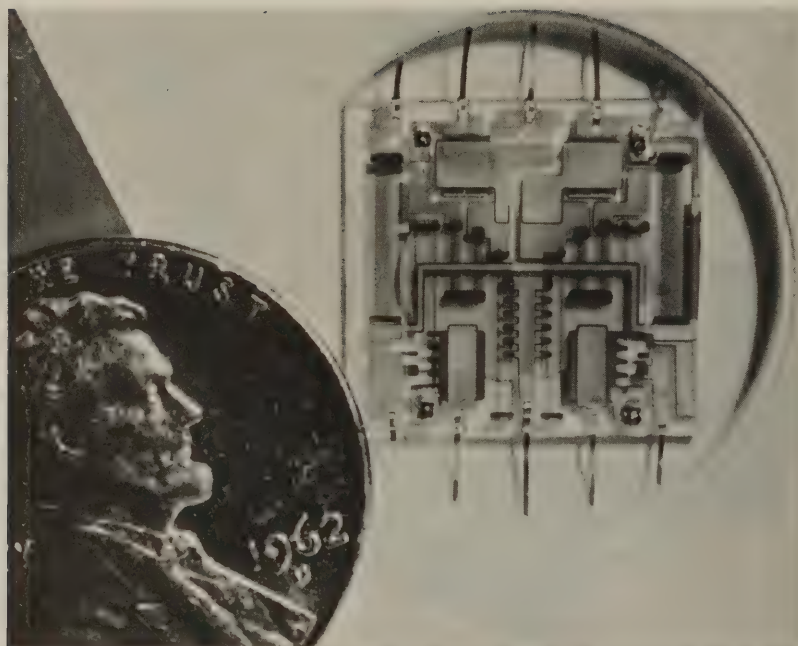
Figure 19

The individual circuit mounted on the header with the leads bonded in place.

Hybrid Integrated Circuits

A HYBRID INTEGRATED CIRCUIT is one in which individual components are in some way ATTACHED to a ceramic or glass substrate, rather than diffused or grown in a single substrate, as with a monolithic circuit. In this respect, hybrid integrated circuits are similar to conventional discrete component circuits. The components on a hybrid circuit may be formed by techniques such as SPUTTERING and VACUUM-DEPOSITION. These are specialized techniques for placing thin layers of materials on a substrate. Circuits made by depositing the resistive and capacitive components in this manner are called THIN-FILM circuits. Some hybrid integrated circuits use thick layers of materials on a substrate to form the resistors. These units are called THICK-FILM circuits. The active devices (transistors and diodes) are attached to the substrate and wired into the circuit before it is packaged. The active devices may be chips or discrete packaged devices.

Multi-chip hybrid circuits may use any combination of thin-film components, monolithic components, or monolithic circuits all interconnected and packaged together. The thin-film components may be individual resistors and capacitors. Since the resistors and capacitors are radically different from conventional passive components, the complete circuit is very small. The transistors and diodes may be in the form of monolithic dice, produced using the monolithic techniques discussed previously. Thus, the multi-chip hybrid circuit uses DIE BONDING to attach these components to a HEADER, or substrate. The active components are interconnected using very light gauge wire (.001 in. diameter). This procedure is called WIRE BONDING.



The above photo shows a thin film circuit with the various active components deposited and bonded to a glass substrate.

Courtesy General Electric Co.

COMPONENT LIMITATIONS

Integrated circuits require special considerations in component characteristics. In addition, these considerations are different for each of the various types of integrated circuits. In monolithic circuits, the component considerations are dependent upon the properties of P-N junctions. Since all the components in a monolithic circuit are formed on a common substrate, they must be electrically isolated from one another. This is accomplished by the isolation diffusion which creates N-type islands in a P-type silicon block as discussed previously. This creates, in effect, a pair of back-to-back diodes between any two islands. Thus, regardless of the polarity between any two components, a back-biased diode is always present between them.

Between the components the isolation diffusion introduces capacitance which must be taken into account. This is because a reverse biased P-N junction forms a double layer which acts as a JUNCTION CAPACITOR. Resistors are also made in monolithic circuits by utilizing the bulk resistivity of the silicon material. Since these properties are inherent in the semiconductor material, the components must be closely controlled in order to obtain only the desired characteristic. In actual practice, however, it is not always possible to completely eliminate one or more of the undesired properties. When an undesired property exists in a component, it is called a PARASITIC. This is

The following Practice Exercise questions cover the subjects which you have just studied. They are:

BASIC CONSTRUCTION

MONOLITHIC INTEGRATED CIRCUITS

CRYSTAL GROWTH

CIRCUIT FABRICATION

HYBRID INTEGRATED CIRCUITS

1. What is the major difference between discrete and integrated circuits.
2. A perfect silicon crystal would be a (a) good conductor, (b) nonconductor.
3. Perfect silicon crystals do not occur in actual practice due to trace amounts of dirt, metal, or other foreign substances called _____.
4. The process of adding impurities to silicon or germanium is called _____.
5. By adding free electrons to a crystal, it is possible to form a P-type semiconductor. True or False?
6. There are two different types of CHARGE CARRIERS in semiconductor conduction. True or False?
7. Polycrystals are materials which have unidirectional lattice structure. True or False?
8. Impurities may be added to a crystal as the crystal is being formed. True or False?
9. The EPITAXIAL GROWTH method of crystal formation takes place in a (a) semiconductor melt, (b) reactor.
10. How do the epitaxial and diffusion processes differ?
11. Through the process of diffusion, it is possible to convert either P or N doped silicon to the opposite type. True or False?
12. What material placed on the wafer permits selective diffusion?
13. When photoresist is exposed to light, does it remain hard or become soft?

- 14. Crystal growth techniques are used mainly to form the circuit substrate while diffusions are used to form components and transistor elements. True or False?**
- 15. How are the various sections on an integrated circuit interconnected?**
- 16. How are the integrated circuits, formed on a wafer, separated?**
- 17. How does a hybrid integrated circuit differ from a monolithic integrated circuit?**

especially true of capacitance in a monolithic circuit. Since the parasitic cannot be eliminated, it must be incorporated or tolerated in the circuit.

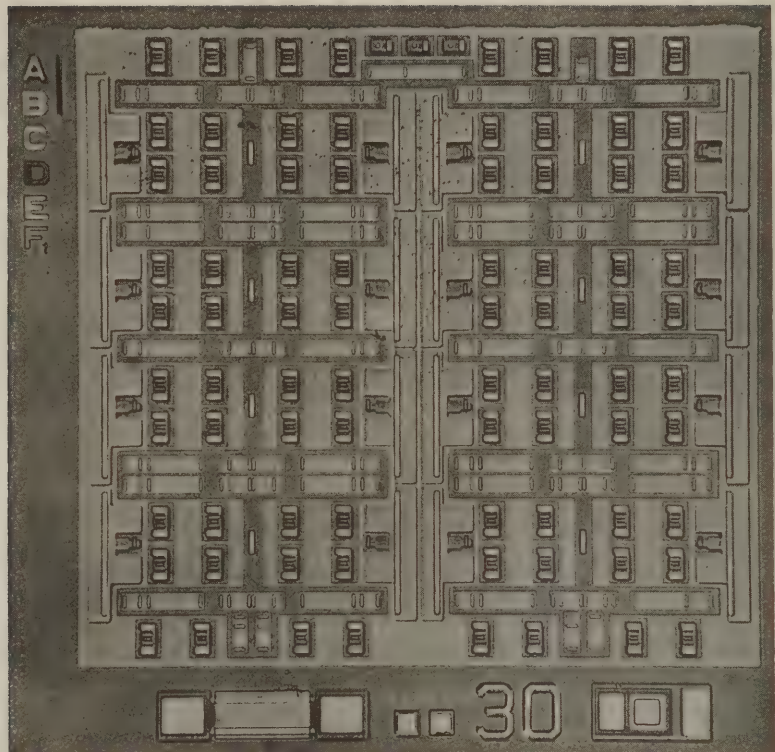
LINEAR CIRCUITS

LINEAR CIRCUITS use active and/or passive components to form analog devices. An analog device or circuit is one in which the output is equal or proportional to the input. The output may be out-of-phase from the input. Because of this analog characteristic, linear circuits are useful as amplifiers and oscillators in radio, TV, and high-fidelity sound systems.

The component limitations discussed previously which apply to integrated circuits apply most particularly to linear circuits in IC (integrated circuit) form. This is because linear circuits often require large values of inductance, capacitance and resistance for their operation. Inductance, at present, is difficult to obtain in IC form, even in small values. Capacitance and resistance values are largely a function of the area, or "real estate", they occupy on a silicon chip. These factors are, of course, inconsistent with the concept of microelectronics, since the aim of microelectronics is to produce smaller devices.

To further illustrate the point, consider the fact that typical IC chip sizes range from 1600 to 3600 square mils (1 mil = 1/1000 inch). A chip of this size will hold several transistors and diodes. However, if this entire area were used to fabricate a single capacitor, the capacitance would be so small that it would be barely usable, even at high frequencies. The same area will hold approximately 200K ohms of resistance, which is not much for the purpose of linear circuits. For these reasons, linear circuits in IC form have been slow to develop and may be more expensive than similar discrete component circuits. However, techniques have been found for producing "basic" linear circuits which fulfill the needs of many linear circuit applications.

Although close-tolerance individual components are very difficult to produce in IC form, accurate resistance ratios and closely matched components are easily obtained. In this respect, discrete component and IC circuits are the opposite of each other. Thus, a matched-component requirement, which is sometimes a disadvantage in discrete component design, represents a distinct advantage in IC design. This capability results from the fact that the IC components are fabricated simultaneously with all conditions being the same at the time of manufacture. Also, the components can be placed in very close proximity to each other, and thus, temperature variations between them are small.



Complex integrated circuits like the one shown above often employ several metallized layers to provide the necessary interconnections.

*Courtesy International Business Machines Corp.,
Components Div.*

Differential Amplifiers

The differential amplifier has become a basic circuit in linear IC technology. It not only uses matched components, but can, in some cases, use active devices in place of space-consuming passive devices. This illustrates another difference between the design philosophy of discrete-component and IC circuits. An increase in the number of active devices in IC circuitry does not necessarily cause a proportional increase in price. Active components require little space and are produced simultaneously. Passive components, however, must be held to a minimum. The reverse is true in discrete component circuitry.

The differential amplifier is a DIRECT-COUPLED amplifier circuit, and further reduces the requirement for using inductors and large-value capacitors. It exhibits excellent stability over a wide frequency range and has a high immunity to interference signals. Thus, for all these reasons, the differential amplifier is an ideal circuit for producing in IC form.

In order to see how all of this is made possible, we will analyze a simple differential amplifier circuit, shown in Figure 20. The matched transistors are each placed in one leg of a bridge circuit. Assuming that the transistors

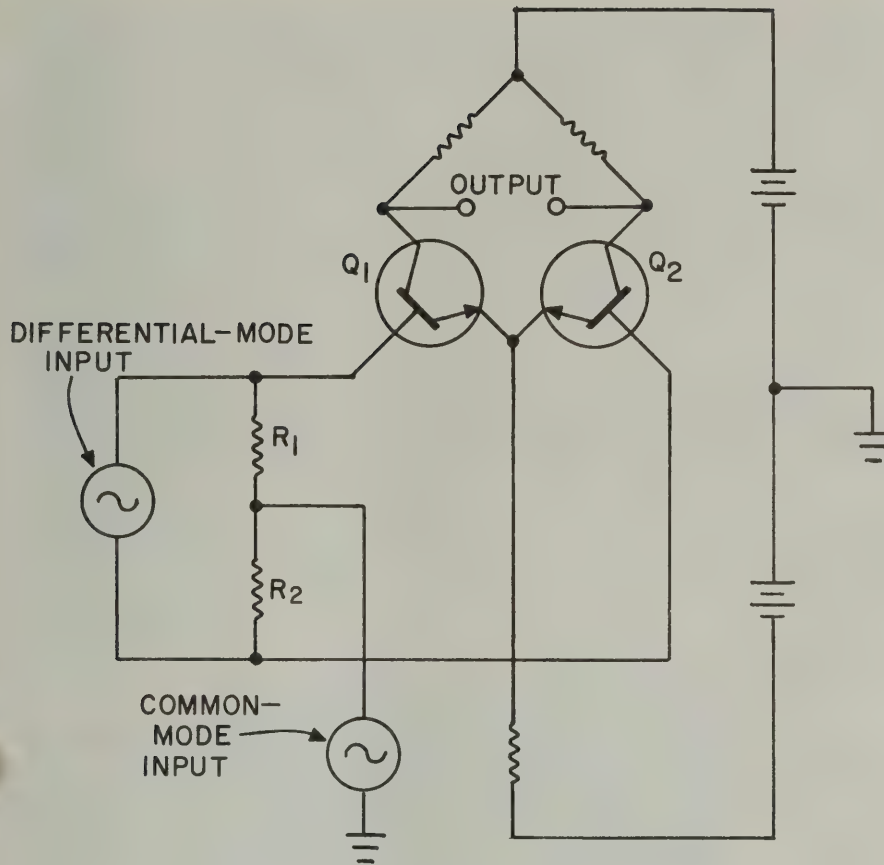


Figure 20

and their collector resistors are identical, and with no signal applied, the bridge would be balanced, and the voltage across the output terminals would be zero. However, with a signal appearing at the DIFFERENTIAL-MODE input (with $R_1 = R_2$), the voltages at the bases of Q_1 and Q_2 are equal in amplitude, but 180° out of phase. Thus, when Q_1 collector current increases, Q_2 collector current decreases the same amount. This produces a difference voltage across the bridge output which is proportional to the gain of the transistors.

A COMMON-MODE signal such as interference, hum, line pickup, or other undesirable noise appearing at the common-mode input is rejected by the amplifier. Since the common-mode signal is applied between the transistor bases and ground, (through matched input resistors), the signals are both equal in amplitude and IN PHASE at both bases. Thus, the current through Q_1 due to the common-mode signal equals the current through Q_2 due to the common-mode signal, and no difference in potential appears between the output terminals. Therefore, we see that the amplifier exhibits high gain for differential-mode inputs, but has zero output for common-mode signals.

The components in a differential amplifier can be CLOSELY matched, but, of course, not perfectly matched. Perfect symmetry does not exist, even though it can be approached in integrated circuitry. Slight mismatches in the components must be taken into account in analyzing a practical differential amplifier circuit. These mismatches affect the operation of the circuit to some degree at both the differential and the common-mode inputs. At the differential-mode input, mismatches in the components cause distortion, or degradation, of the desired output. However, mismatches cause a potentially more serious problem with common-mode signals, since this produces an output where none is desired. A well-balanced differential amplifier produces only a very small DIFFERENCE signal, or output from the common-mode input. The ability of a differential amplifier to prevent the conversion of a common-mode signal into a difference signal at the output is called the COMMON-MODE REJECTION (CMR) RATIO. The CMR ratio is found by first dividing the output difference voltage by the differential input voltage. The resultant is then divided into the common-mode input voltage. The final result is expressed in dB. A well-designed integrated circuit differential amplifier has a CMR ratio of from 100 to 120 dB.

Operational Amplifiers

The term "operational amplifier" stems from the early analog function of these devices in performing mathematical operations, i.e., addition, multiplication, integration, etc. Although operational amplifiers still perform these functions, the term today is defined in a somewhat broader sense. A description such as "a high-quality, high-gain, direct-coupled amplifier assembly, often incorporating provision for overall feedback", is frequently given as the definition for an operational amplifier. Originally, the definition was much more specific, since operational amplifiers were used almost exclusively in analog computer systems. However, with the advent of integrated circuit technology, these high-quality amplifiers can be produced at much lower cost. As a result, many more applications for operational amplifiers now exist which previously did not, since the expense is no longer prohibitive.

An operational amplifier is not, by itself, a complete piece of equipment. It is designed to be inserted into other, larger equipment. Because of this, the

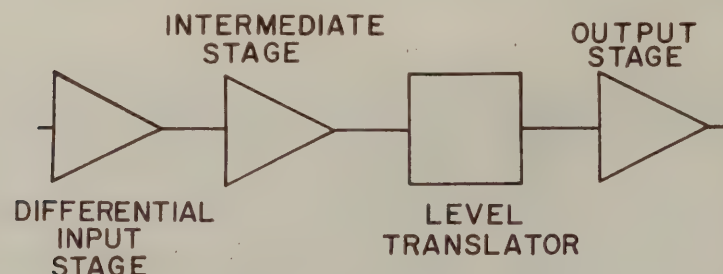


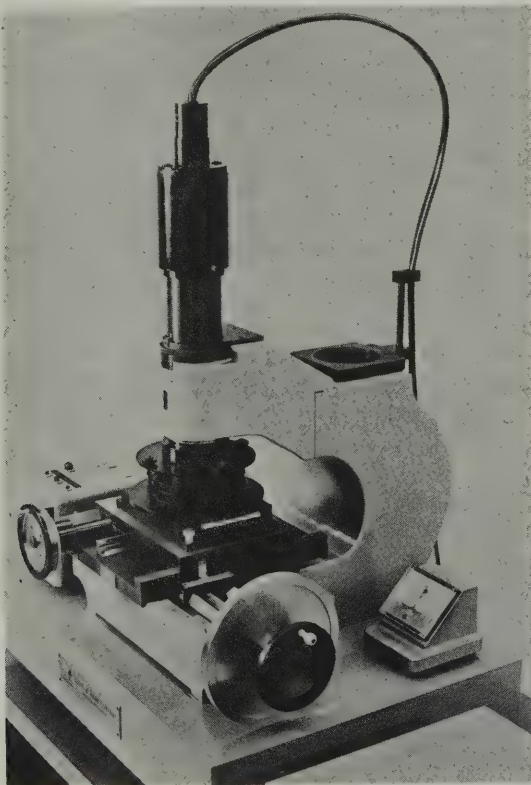
Figure 21

aim of integrated circuit designers is to make operational amplifiers as flexible and universal as possible. Therefore, a standard block diagram format has been adopted by nearly every operational amplifier ("Op-Amp") manufacturer, as shown in Figure 21.

The first stage of an op-amp is a differential amplifier which provides most of the circuit gain. This stage should be well designed and have a high CMR ratio. The design of the first stage is important because imperfections here will affect succeeding stages, and thus the output.

The second stage is provided mainly for additional circuit gain. The input resistance should be high, to prevent loading of the first stage. The second stage may be another differential amplifier or a conventional amplifier. An emitter-follower configuration is often employed.

In a direct-coupled system such as this, a dc voltage level is propagated through the chain of stages. Thus, the amplifier output would have a dc component in addition to the desired ac output signal. In order to provide an output signal which varies about a zero reference level, a level translator



The photorepeater shown above is a precision machine that is used to expose the photosensitive material used to create the diffusion masks. It has a step-and-repeat capability so that many circuits can be formed on one wafer.

Courtesy David W. Mann Co.

stage is included. This is accomplished using sum-and-difference bias techniques (common in the analog art) in this stage.

The output stage, or final amplifier, generally consists of a signal driver and an output transistor pair. The final stage employs negative feedback from the amplifier output to the final driver input. This feedback, common in all high-quality amplifiers, compensates for nonlinearities in the output stages, and assures a symmetrical output signal.

DIGITAL CIRCUITS

Digital circuits gained the advantage of the early research in integrated circuit development. The demands of the computer industry for smaller, cheaper, and higher speed switching circuits, in addition to the simplicity of the circuits themselves, were chiefly responsible for the early lead established in digital integrated circuitry.

A digital circuit is basically a switch. In computers and other digital devices, some of the switches are called GATES. These gate circuits perform a logical decision-making function, and are called gate LOGIC circuits. Three of the major logic lines, or families, are resistor-transistor logic (RTL), diode-transistor logic (DTL) and transistor-transistor logic (TTL).

Most logic circuits employ what is called SATURATED-MODE switching. This means that the transistors used are driven from deep cutoff to high level conduction (saturation). The high conduction results in a very low collector voltage, which remains relatively constant even when the base forward bias voltage is increased.

Three logic functions performed by logic gates are AND, OR, and NOT. An AND gate or an OR gate may be coupled to a NOT gate. In this case, the two gates together are called NAND and NOR gates, respectively.

Resistor-Transistor Logic (RTL)

Prior to the widespread use of integrated circuits, **RESISTOR-TRANSISTOR LOGIC (RTL)** was the most commonly used form of digital logic circuit. RTL circuits use relatively few active components, and are therefore most suitable for use in discrete component circuits. Since logic designers were most familiar with this type, it was the first to be produced in integrated circuit form. In fact, the earliest integrated circuit designs were direct adaptations from discrete component designs.

A basic three-input RTL gate circuit is shown in Figure 22A. In this circuit, transistors are used as active element switches which have only two basic levels of operation. These are often designated as "on" and "off" or "high" and "low". These terms generally refer to output voltage levels and not necessarily the current flow through the device. Remember that in saturated

switching, the "high" voltage level occurs when the current through the device is the smallest.

In an RTL circuit, a resistor is placed in each of the inputs to ensure more uniform conduction of each transistor. This raises the input resistances and minimizes the differences between the transistors. Thus, the inputs are more closely matched to each other. The resistors improve the circuit operation, but only at the expense of switching speed and/or power dissipation. (It is inherent in logic circuits that the faster the speed, the higher the internal power dissipation in a given circuit design).

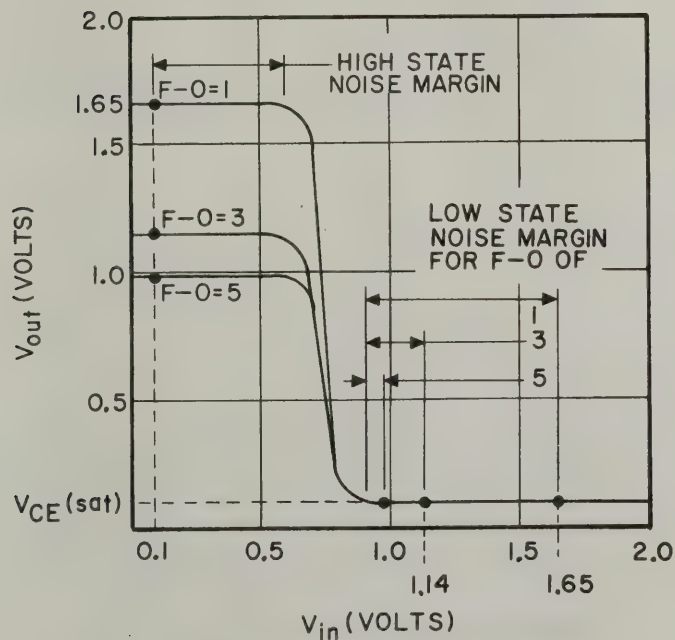
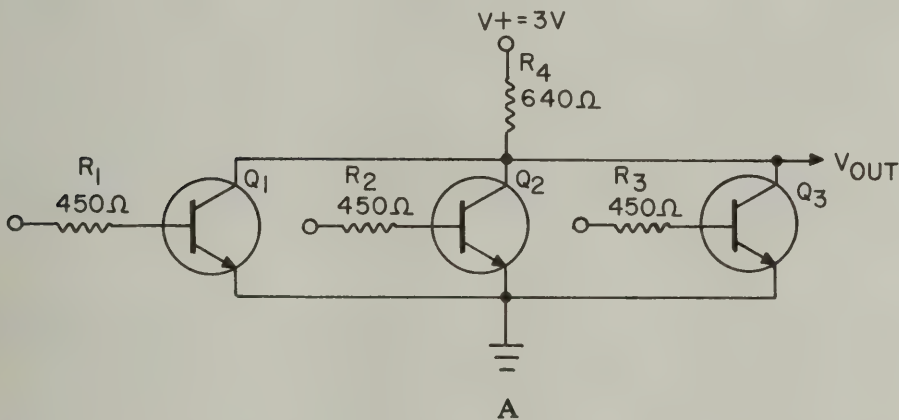


Figure 22

The power dissipation values in RTL circuits can be varied within narrow limits by changing the value of the resistors. Most commercial circuits are available in two power dissipation ranges. Consequently, the speed is also affected.

The characteristic curves for an RTL gate, shown in Figure 22B, illustrate some of the other important properties of RTL circuits. One of these is FAN-OUT (F-O). Fan-out is an important specification in digital circuits, and refers to the number of LOADS connected to the gate. In other words, the fan-out is the number of circuits the output is required to drive. The figure shows the output voltage on the vertical axis at the left. Note that for the "high" state, the output voltage is maximum, or nearer 2V, and the transistor is cut off. For the "low" state, the transistor is conducting at saturation, nearer (or at) $V_{CE}(SAT.)$. The transition region (threshold) occurs on the horizontal axis (V_{IN}) near 0.7 volts. As the voltage to one of the inputs rises to approximately 0.7 volts, the transistor begins to conduct, and the gate output drops rapidly toward $V_{CE}(SAT)$ (approx. 0.1 volt). For a fan-out of 1, the output voltage may vary over a range of 0.1 volt (conducting), to 1.65 volts (cutoff). For a fan-out of 3, the range is from 0.1 to only 1.14 volts. For a fan-out of 5, the range is still smaller—0.1 to 0.98 volts.

Looking at the curve for a fan-out of 1, note what happens when NOISE reaches the gate inputs. When the inputs are "low" (0V), the output is at 1.65V. Notice that a noise pulse of 0.5V or slightly higher would drive the transistor toward the transition region, and the gate would begin to change states ("high" to "low"). However, we can say that the gate is immune to noise signals UP TO 0.6V, or it has a HIGH-STATE NOISE MARGIN of 0.6 volt. It is also important to note that the high-state noise margin is the same for all three fan-outs (1, 3, and 5). As a general statement, we may say that THE HIGH-STATE NOISE MARGIN IS INDEPENDENT OF FAN-OUT.

This is not the case for the low-state noise margin. Assume a "low" (output) condition exists with the gate INPUT (high) at 1.65 volts (F-O of 1). In this case, a NEGATIVE noise spike of 0.75 volt would drive the transistor back toward the transition region, and cause the gate to change states from low to high. If the gate were operating at an input voltage of 1.14 volts (F-O of 3), a negative pulse of 0.2 volt would cause the gate to change state. For a fan-out of 5, a noise pulse of only 0.08 volt would cause the gate to change states. Therefore, we see that the LOW-STATE NOISE MARGIN IS AFFECTED TO A CONSIDERABLE DEGREE BY FAN-OUT. In an RTL gate circuit, the fan-out is somewhat limited, and the noise margin becomes small as the fan-out is increased. Nonetheless, RTL circuits are widely used and remain a popular form of integrated logic.

Diode-Transistor Logic (DTL)

DIODE-TRANSISTOR LOGIC (DTL) has performance advantages over RTL, but, due to the increased number of active components, is more ex-

pensive in discrete-component circuits. As we have seen, however, an increase in active components in integrated form does not necessarily produce a proportional increase in cost. In addition, the diodes in DTL, which replace resistors in RTL, require less space on an integrated circuit chip.

A DTL circuit, shown in Figure 23A, is characterized by the use of diodes at the inputs (D_1 and D_2) to inject the input signals. One or two other diodes (D_3 and D_4) are also used in series with the base of the transistor. The

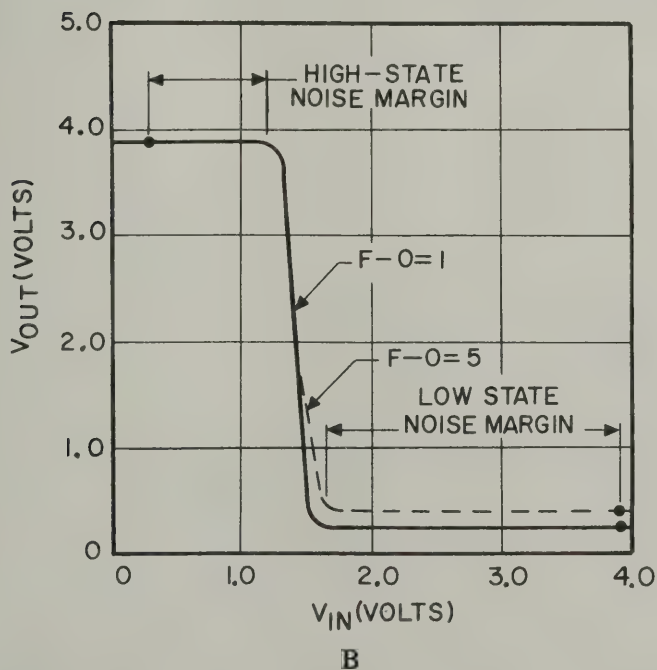
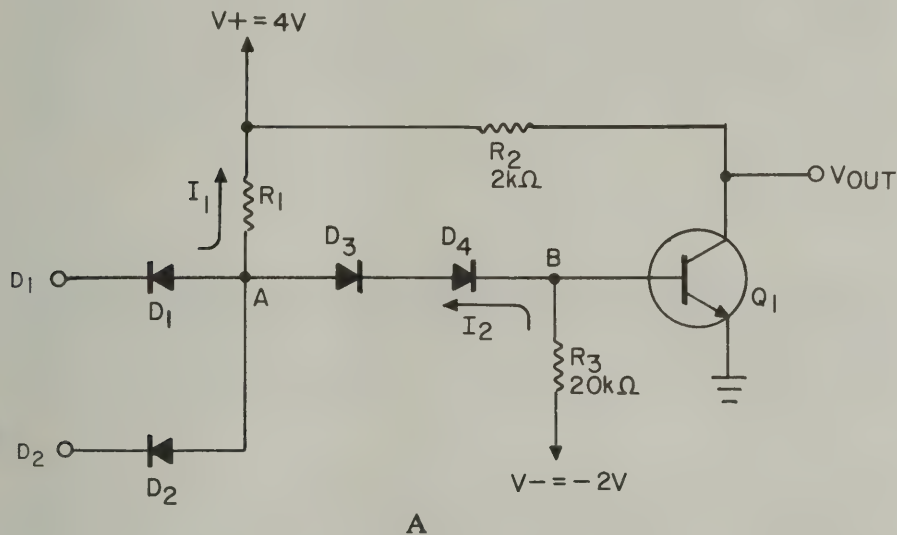


Figure 23

voltage drop across the two series diodes, D_3 and D_4 , provides bias for Q_1 . The diodes also stabilize the current through the transistor in much the same way as do the resistors in an RTL gate.

To examine circuit operation, suppose a low state voltage, $V_{CE(sat)} = 0.1$ volt, is applied to the D_1 input. This forward biases D_1 , resulting in the current I_1 shown through R_1 . With D_1 forward biased, the voltage at point A will be the voltage drop across the diode, approximately 0.7 volts. With point A at 0.7 volt, D_3 and D_4 will be forward biased by V_- and current I_2 will flow as shown. The voltage at point B will be the drop across D_3 and D_4 minus the voltage at point A. Since the voltage at point A is 0.7 volts and the drop across D_3 and D_4 is 0.7 volts each, the voltage at point B will be -0.7 volts. This reverse biases Q_1 and the resultant collector voltage, V_{OUT} , is equal to V_+ or 4 volts. The same action occurs if a low state voltage is applied to the D_2 input.

To change the state of the circuit, a high state voltage must be applied to BOTH the D_1 and D_2 inputs. This causes both D_1 and D_2 to be reverse biased. As a result, the $+V$ supply forward biases Q_1 through R_1 . With Q_1 saturated, its collector voltage drops to the low state, $V_{CE(sat)} = 0.1$ volt.

Since the gating function is controlled by diodes in the DTL circuit, the high and low state voltage levels are not affected by fan-in or fan-out to the extent that occurs in RTL circuits. The result is a greater noise margin for

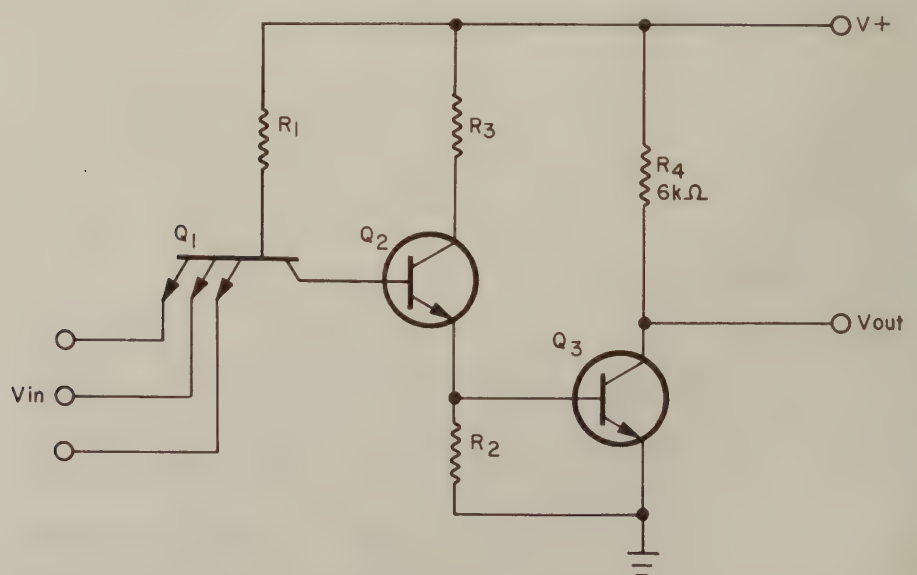


Figure 24

DTL circuits. Figure 23B shows the characteristic curves for the DTL circuit. Note that the low state noise margin is basically unaffected by fan-out. Compare this to the curves for the RTL circuit shown in Figure 22B.

Transistor-Transistor Logic (TTL)

One of the fastest saturated logic systems in use today is TRANSISTOR-TRANSISTOR LOGIC (TTL), or as it is often called, T²L. TTL differs from RTL and DTL in that transistors are used as the input gating elements. In most cases, this is accomplished by simply employing an input transistor with multiple emitters.

Figure 24 shows the circuit of a typical TTL logic gate. When one or more of the emitters of Q_1 is at the low state, the emitter junction of Q_1 is forward biased. With Q_1 forward biased, the voltage at its collector is low and not high enough to drive Q_2 and Q_3 into conduction. As a result, Q_3 is cut off, and its collector voltage, V_{OUT} , is at the high state, V_+ . When all the inputs of Q_1 are at the high state, V_+ , the Q_1 emitter junction is reverse biased. As a result, the Q_1 collector voltage is high enough to forward bias Q_2 , and in turn forward bias Q_3 . With Q_3 forward biased, its collector voltage is low and the output voltage V_{OUT} drops to $V_{CE(sat)} = 0.1$ volts.

The fan-out and noise margin characteristics of the TTL circuits are similar to those of the DTL circuits. The principal advantage of TTL is its high operating speed.

SUMMARY

Integrated circuits differ from conventional or discrete circuits in that all the components of an integrated circuit are in one package. Discrete circuits employ separately packaged components such as resistors, diodes, capacitors and transistors.

The most popular type of integrated circuit is constructed in monolithic form. In this type of unit, transistors, diodes, resistors and capacitors are formed on a single piece of silicon material called a substrate. The various P-type and N-type layers are formed by epitaxial and diffusion techniques. Photo etching techniques are used to control the areas where P and N-type regions are formed, to construct diodes, transistors and other components.

Integrated circuits are also constructed in thin-film and thick film form. With this technique, resistors and capacitors are made on a glass or ceramic substrate, individual transistors and diodes are bonded to the passive components and the finished circuit is then placed in a single package.

Integrated circuits are generally classed as either linear or digital. Linear circuits generally take on the form of oscillators, amplifiers, regulators and similar devices. Differential amplifiers are widely used in linear integrated circuits. The differential amplifier takes advantage of the component matching available in integrated circuits to provide good noise rejection and high gain.

Digital integrated circuits are those employed to perform various logic functions such as OR, AND, NAND and NOR. There are several different types of basic logic circuits. Resistor-transistor logic circuits, RTL, employ resistor inputs to the logic function stages and the outputs are usually followed by an inverting transistor amplifier. RTL circuits suffer from poor noise immunity and slow speed operation. Diode transistor logic circuits, DTL, employ diode inputs to the logic function circuits. A transistor amplifies and inverts the output. DTL circuits have a better noise margin and higher operating speed than RTL circuits. Transistor-transistor logic circuits, TTL, have the highest operating speed. Transistors are used in this arrangement to perform the logic function as well as to amplify and invert the output.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

COMPONENT LIMITATIONS

LINEAR CIRCUITS

DIFFERENTIAL AMPLIFIERS

OPERATIONAL AMPLIFIERS

DIGITAL CIRCUITS

RESISTOR-TRANSISTOR LOGIC (RTL)

DIODE-TRANSISTOR LOGIC (DTL)

TRANSISTOR-TRANSISTOR LOGIC (TTL)

18. What is a parasitic in an integrated circuit?
19. Which generally occupies a greater area on an integrated circuit, a resistor or a transistor?
20. Matched components are easy to obtain with integrated circuits. True or False?
21. What is the name given to the amplifier circuit shown in Figure 20?
22. Why does a differential amplifier reject hum and noise signals?
23. A positive signal at the base of Q_2 and a negative signal at the base of Q_1 in Figure 20 will cause the Q_2 collector voltage to increase. True or False?
24. What name is given to the circuit arrangement shown in Figure 21?
25. What amplifier configuration can be used to provide a high input impedance and thus minimum loading?
26. What are the two conditions of a transistor when it is operated as a saturated mode switch?
27. A circuit which performs a logic function such as NAND or AND is generally called a _____.
28. In the circuit of Figure 22A, if the input to only Q_1 is high, and the Q_2 and Q_3 inputs are low, the output is low. True or False?

29. How is a high output obtained in the circuit of Figure 22?
30. Is the noise margin of an RTL circuit increased or decreased as the fan-out is increased?
31. When the input to D_1 is low and the input to D_2 is high in Figure 23, the voltage at point A is low. True or False?
32. With the conditions mentioned above in Exercise 31, is D_2 conducting or cut off?
33. The noise margin of a T^2L circuit is less than that of an RTL circuit. True or False?

IMPORTANT DEFINITIONS

DIFFUSION — A thermal process by which P or N-type regions can be formed in a substrate material.

DIGITAL CIRCUIT — A circuit whose inputs determine which of the two possible output states the output assumes.

DIODE-TRANSISTOR LOGIC (DTL) — A type of logic in which the logic function is performed by a diode network. A transistor is used to amplify and invert the output.

DISCRETE COMPONENT — A separate, individually constructed and packaged component such as a resistor, diode or transistor.

EPITAXIAL GROWTH — A process by which P or N-type material is grown on a substrate material.

FAN-OUT — The number of circuits connected to the output of a gate circuit.

HYBRID INTEGRATED CIRCUIT — An integrated circuit made up of individual components or groups of components that are attached to a glass or ceramic substrate.

INTEGRATED CIRCUIT (IC) — A small electrical device built into a single package containing several or many components which may perform a complete circuit function.

LINEAR CIRCUIT — A circuit in which the output follows the input in a linear or analog manner.

MONOLITHIC INTEGRATED CIRCUIT — An integrated circuit constructed on a single piece of substrate material.

RESISTOR-TRANSISTOR LOGIC (RTL) — A logic circuit in which the logic function is performed by a resistor network. A transistor is used to amplify and invert the output.

SUBSTRATE — The base material upon which the components of an integrated circuit are constructed.

TRANSISTOR-TRANSISTOR LOGIC (TTL OR T²L) — A type of logic circuit in which the logic function is performed by transistors, generally with multiple emitters. Another transistor is used to amplify and invert the output.

PRACTICE EXERCISE SOLUTIONS

1. An integrated circuit can be a complete circuit function in one case. A discrete circuit is made up of individually packaged circuit components such as resistors, transistors, diodes etc.
2. (b) nonconductor — there are no free current carriers in pure silicon material.
3. impurities
4. doping
5. False — P-type material is formed by adding an acceptor impurity which causes the loss of an electron or the generation of a hole.
6. True — Both holes and electrons serve as charge carriers in semiconductor material.
7. False — Monocrystals have a unidirectional lattice structure.
8. True — The impurities can be added as the crystal is grown from a seed or grown by the epitaxial method.
9. (b) reactor
10. In the epitaxial process a new layer of P or N-type material is grown on the substrate. The diffusion process causes a P or N-type region to be formed in the existing substrate.
11. True — The impurity concentration is increased until it overcomes the impurity concentration of the other type of material.
12. The silicon dioxide layer.
13. Generally, the photoresist used in integrated circuit work becomes hard when exposed to light.
14. True
15. A metallized layer is used to interconnect the various sections of an integrated circuit.
16. The integrated circuits are separated by scribing the wafer with a diamond tipped tool. This breaks the wafer into individual dice.
17. The hybrid integrated circuit consists of individual or combinations of components mounted on a glass or ceramic substrate. In the monolithic integrated circuit, all components are formed on the same silicon substrate at the same time.

PRACTICE EXERCISE SOLUTIONS (Continued)

18. A parasitic is an undesired component, such as a capacitor formed in an integrated circuit.
19. A resistor (or capacitor) generally occupies more area than a transistor.
20. True — Since the components are made at the same time they can be easily matched.
21. The circuit is called a differential amplifier.
22. Hum and noise signals are generally common mode signals. They receive equal amplification and cancel at the differential output.
23. False — Since Q_1 and Q_2 are NPN transistors, the positive signal at the Q_2 base will increase collector current and decrease collector voltage.
24. The circuit of Figure 21 is an IC operational amplifier.
25. An emitter follower configuration produces minimum loading since it has a high input impedance.
26. When operated as a saturated mode switch, a transistor operates either in cutoff or saturation.
27. gate.
28. True — Q_1 saturates and drives the output down to the low level, $V_{CE(sat)} = 0.1$ volts.
29. The output is high only when all inputs are low.
30. The noise margin decreases as the fan-out increases.
31. True — the voltage at point A is low, about 0.7 volts.
32. D_2 is cut off, since the voltage at its input reverse biases D_2 .
33. False



ONE OF THE

BELL & HOWELL SCHOOLS

EXAMINATION
CHECK SHEET

III2A

1. B - THE BASE MATERIAL THAT AN INTEGRATED CIRCUIT IS CONSTRUCTED UPON IS REFERRED TO AS THE -- SUBSTRATE.

The base material is generally referred to as the wafer or substrate.

2. A - THE PROCESS OF GROWING A LAYER OF DOPED SEMICONDUCTOR MATERIAL ON A SLICE OF SILICON IS CALLED THE -- EPITAXIAL PROCESS.

In the epitaxial process, silicon gas introduced into a reactor grows a doped silicon layer on the wafer material. The impurities in the gas determine if the grown layer is N or P-type.

3. C - THE LAYER OF MATERIAL DEPOSITED ON AN INTEGRATED CIRCUIT WAFER THAT PREVENTS DIFFUSION UNDER IT IS THE -- SiO_2 LAYER.

The photoetching technique is employed to form the desired pattern for removal of the SiO_2 layer.

4. A - THE VARIOUS AREAS ON THE WAFER THAT MAKE UP AN INTEGRATED CIRCUIT ARE INTER-CONNECTED WITH A -- METALLIZED LAYER.

The metallized layer is etched to form the desired interconnection pattern.

1.

5. D - THE COMPONENT WHICH OCCUPIES THE MOST SPACE ON AN INTEGRATED CIRCUIT IS A -- CAPACITOR.

Resistors and capacitors require large areas and as a result their values are limited.

2.

6. A - THE TYPE OF AMPLIFIER CIRCUIT MOST OFTEN USED IN LINEAR INTEGRATED CIRCUITS IS THE -- DIFFERENTIAL AMPLIFIER.

The differential amplifier requires matched components which are easy to obtain in integrated circuit form.

3.

7. A - IN THE CIRCUIT OF FIGURE 20, IF THE INPUT SIGNAL DRIVES THE BASE OF Q_1 POSITIVE AND THE Q_2 BASE NEGATIVE -- Q_1 COLLECTOR VOLTAGE DECREASES.

The positive signal at the Q_1 base increases forward bias, increasing Q_1 collector current and decreasing collector voltage.

4.

8. B - COMPARING RTL TO DTL CIRCUITS, -- DTL CIRCUITS CAN SWITCH FASTER THAN RTL.

Both DTL and RTL circuits operate as saturated switches.

5.

9. B - IN A DTL CIRCUIT, THE LOGIC FUNCTION IS PERFORMED BY -- DIODES.

The diodes perform the logic function and the transistor amplifies the output of the diode logic circuit.

6.

10. A - OF THE FOLLOWING CIRCUIT TYPES, WHICH HAS THE POOREST NOISE MARGIN? -- RTL.

DTL and TTL circuits have a greater noise margin since diodes in DTL and transistors in TTL perform the logic function.

7.

8.

9.

10.

11.

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1112A

Example: In our solar system, Pluto is a
 A ☐ B ☒ C ☐ D ☐
 A. star. B. planet. C. dog. D. meteor.

1. A ☐ B ☒ C ☐ D ☐ **The base material that an integrated circuit is constructed upon is referred to as the**
 (A) crystal. (B) substrate. (C) polycrystal. (D) monocrystal.

2. A ☐ B ☒ C ☐ D ☐ **The process of growing a layer of doped semiconductor material on a slice of silicon is called the**
 (A) epitaxial process. (B) diffusion process. (C) etching process. (D) seed pulling process.

3. A ☐ B ☒ C ☐ D ☐ **The protective layer of material grown in an oxygen atmosphere on an integrated circuit wafer that prevents diffusion under it is the**
 (A) photomask. (B) photoresist. (C) SiO₂ layer. (D) final mask.

4. A ☐ B ☒ C ☐ D ☐ **The various areas on the wafer that make up an integrated circuit are interconnected with a**
 (A) metallized layer. (B) SiO₂ layer. (C) epitaxial layer. (D) mask.

5. A ☐ B ☐ C ☒ D ☐ **The component which occupies the most space on an integrated circuit is a**
 (A) diode. (B) transistor. (C) interconnecting lead. (D) capacitor.

6. A ☒ B ☐ C ☐ D ☐ **The type of amplifier circuit most often used in linear integrated circuits is the**
 (A) differential amplifier. (B) parallel amplifier. (C) common-base configuration. (D) push-pull amplifier.

7. A ☒ B ☐ C ☐ D ☐ **In the circuit of Figure 20, if the input signal drives the base of Q₁ positive and the Q₂ base negative**
 (A) Q₁ collector voltage decreases. (B) Q₁ collector current decreases. (C) Q₂ collector current increases. (D) Q₂ collector voltage decreases.

8. A ☐ B ☒ C ☐ D ☐ **Comparing RTL to DTL circuits,**
 (A) RTL circuits can switch faster than DTL. (B) DTL circuits can switch faster than RTL. (C) RTL circuits operate as saturated switches, DTL do not. (D) DTL circuits operate as saturated switches, RTL do not.

9. A ☐ B ☒ C ☐ D ☐ **In a DTL circuit, the logic function is performed by**
 (A) resistors. (B) diodes. (C) transistors. (D) capacitors.

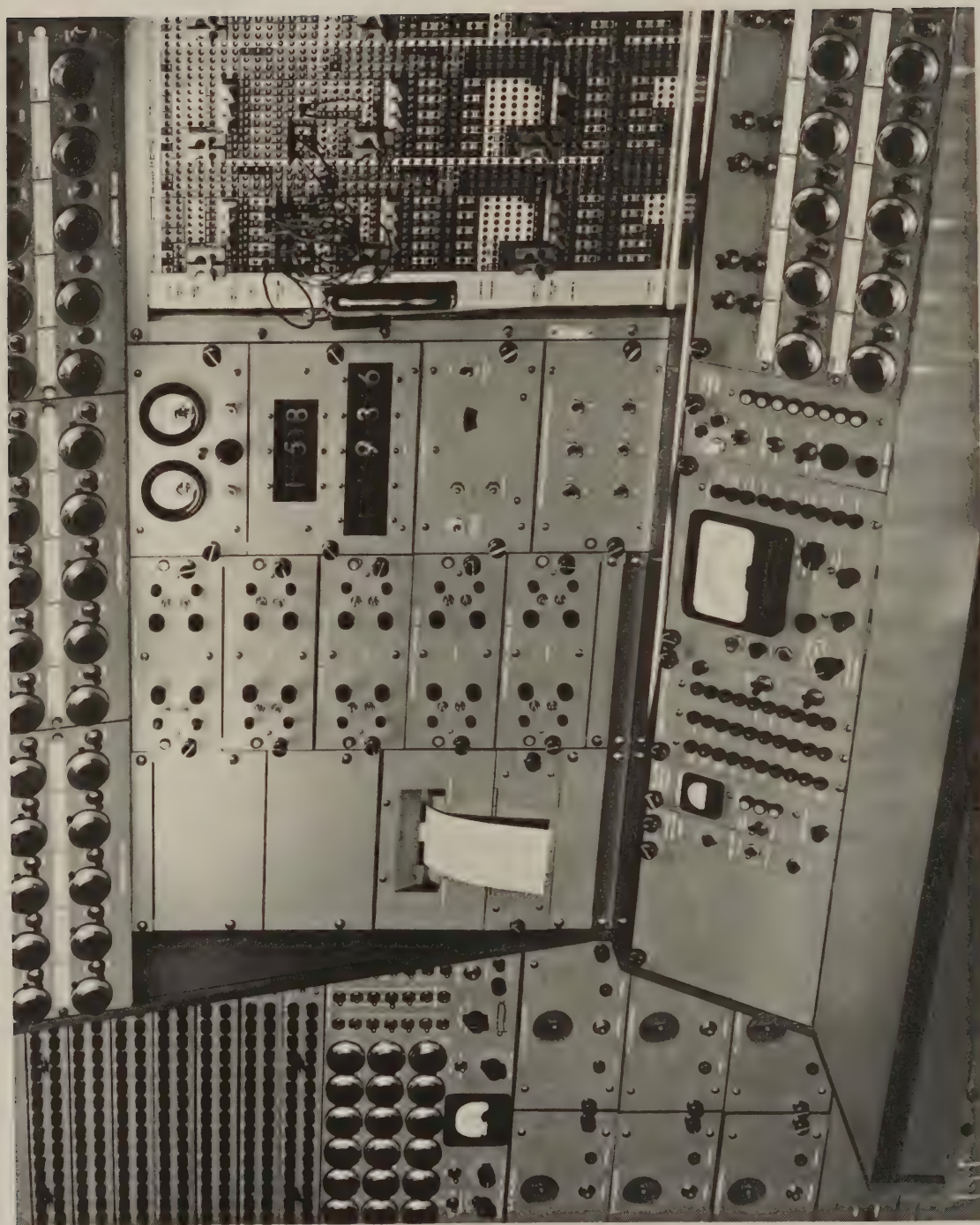
10. A ☒ B ☐ C ☐ D ☐ **Of the following circuit types, which has the poorest noise margin?**
 (A) RTL. (B) DTL. (C) TTL. (D) T²L.

ELECTROMECHANICAL DEVICES

1115

COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.





Electromechanical devices provide the interface means between man and machine. In order to operate this analog computer, the programmer must come into physical contact with the equipment. Some of the electromechanical devices used here allow the programmer to "enter information" into the machine, and others provide a display or readout of the results.

Courtesy Electronic Associates, Inc.

CONTENTS

Switches	Page 4
Leaf Switches	Page 5
Snap-acting Switches	Page 7
Rotary Wafer Switches	Page 9
Relays	Page 11
General Purpose Relays	Page 15
Telephone-type Relays	Page 16
Resonant Reed Relays	Page 22
Magnetic Reed Relays	Page 23
Stepping Switches	Page 26
Solenoids	Page 29
Electric Clutches and Brakes	Page 32

*Things don't turn up in this world until somebody
turns them up.*

—Garfield

ELECTROMECHANICAL DEVICES

Electromechanical devices may be described as a general class of instruments which permit the control or use of electric energy—through mechanical motion. Electrical circuits are capable of performing a wide variety of functions such as heating buildings, providing light and performing arithmetic calculations. Before any of this can be done, however, there must be some means for controlling and using electricity. Since electrical currents must travel through some substance, such as copper wire, one might assume that the control of electricity involves methods of controlling and manipulating wires. If the wires involved are part of an electric circuit, these methods would be controlling electric circuits. Such is indeed the case. The devices used for human control or use of electric circuits are called **ELECTROMECHANICAL DEVICES**.



Figure
1

Some electromechanical devices convert energy from one form to another, and are properly classed as **TRANSDUCERS**. However, this is often an incidental characteristic of the device itself, and is not the primary function of the device. Many electromechanical devices are used to **TRANSFER** energy, not convert it. In this lesson we will consider several devices which transfer energy from one electric circuit to another.

SWITCHES

The simplest type of electromechanical device is a manually operated switch. Through the use of switches, one is able to control the flow of electric currents. This may involve either the completion or interruption of a simple series circuit, or a complex transfer of many circuits within a sophisticated electronic system.

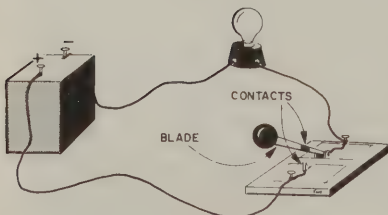


Figure
2

Figure 1 illustrates a very simple switch. As shown, the switch may be nothing more than two wires just touching at some point. This switch would not be very practical for actual use, however, since the wires could be easily jarred loose, causing intermittent operation and poor circuit performance. The idea behind this switch is used more conveniently in the arrangement of Figure 2. Note that the switch merely completes or interrupts the circuit by engaging or disengaging a set of **CONTACTS**. Switches such as this, although simple, are used in power switching applications, and are called **KNIFE SWITCHES**. Switches are classed as to function by the number of **POLES** and **THROWS** in the switch. The knife switch shown is a single-pole, single throw (SPST) switch because it has a single electrical circuit, or **POLE**, and a single closed position (**THROW**). Some knife switches have two closed positions, and they are known as double-throw switches. The poles and throws of a switch are but one way of classifying these devices.

Switches are also classed by their mechanical construction. Typical switches, which we will discuss in this lesson, are leaf switches, snap-acting switches, and rotary wafer switches. These, and other, types are available simply because they are more suitable in one application than another. For example, a rotary wafer switch with multiple contacts and wafers can perform switching actions by one movement that a snap-acting switch would be unable to duplicate.

Leaf Switches

Switches using the principles of contacting wires or blades which move together or apart in response to the application or release of pressure are called **LEAF SWITCHES**. A common application for these switches is in telephone systems. Therefore, these switches are also referred to as telephone-type switches. Figure 3 shows the component parts of a leaf switch.

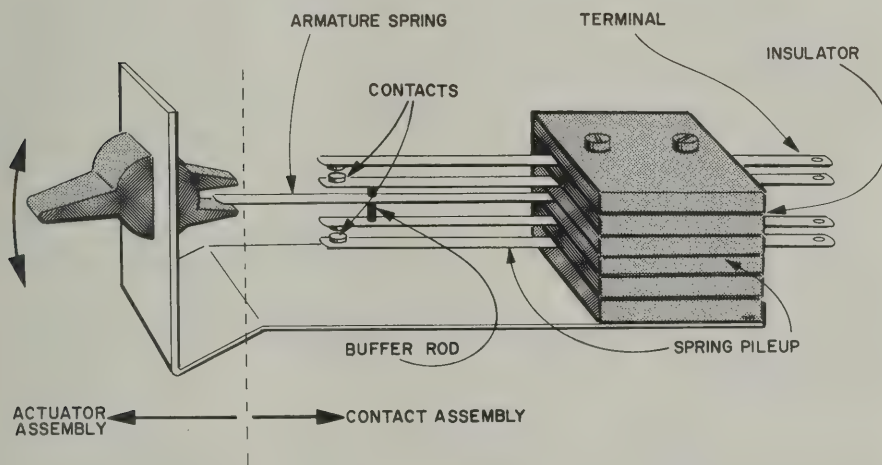


Figure 3

Every switch contains two parts: the CONTACT ASSEMBLY, which makes, breaks or transfers a circuit; and the ACTUATOR ASSEMBLY, which operates the contact assembly. As shown in Figure 4, there are several component parts within the contact assembly. The contact assembly in a leaf switch is called the **SPRING PILEUP**. The spring pileup consists of a number of flat, flexible blades made of a spring-metal alloy. Thus, the blades return to their original flat shape after being stressed or bent, and are properly called **SPRINGS**. Attached to the springs at one end are small metal contacts placed in such a way that the contact surfaces on adjacent springs will touch when the springs are brought together. In this way, a simple switch is formed. The springs on which specific sets of contacts actually touch during a switching operation are called **SPRING COMBINATIONS**. Figure 4 illustrates a leaf

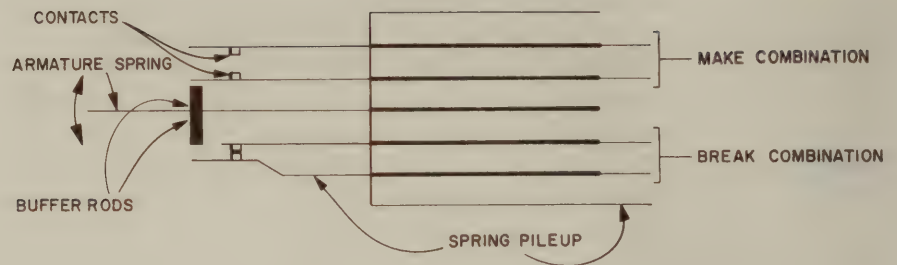


Figure 4

switch that includes two sets of spring combinations: a MAKE-COMBINATION, and a BREAK COMBINATION. In this illustration the center spring is not part of either combination, but merely operates the others, and is called the ARMATURE SPRING. An armature spring is one which is operated or mechanically moved by the switch actuator assembly. Note that as the armature spring is moved upwards, a BUFFER ROD pushes against one of the springs in the make combination. Continued upward motion of the armature spring causes the make combination to close, thus completing a switching circuit. Although the armature spring must travel upward far enough to close the contacts, there is a slight OVERTRAVEL which causes the uppermost spring to bend also slightly upward. This holds the contacts together under spring tension and assures that small vibration will not jar and separate the contacts, causing intermittent connections and poor performance. The contacts in the make combination are said to be normally open (N. O.). That is, the contacts are open when pressure is not being applied to the switch handle. The reverse is true for the contacts in the break combination. In this case, the contacts are normally closed (N. C.), whenever the switch handle is not being operated.

The number of spring combinations in a leaf switch need not be limited to two as shown in Figures 4 and 5. The practical limit for the number of spring combinations is determined by the opposing force offered by a large number of springs. As the spring pileup becomes large, the switch becomes stiff and difficult to operate comfortably.

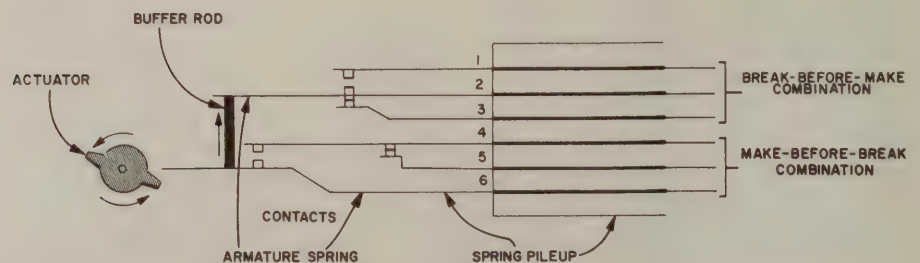
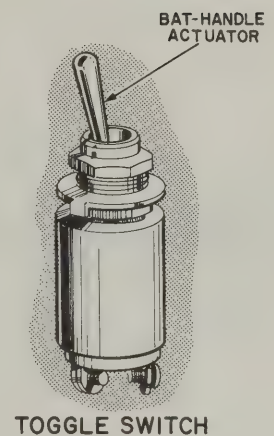


Figure 5

Leaf switches can be designed to provide more complex functions than on-off switching. In addition to the make and break combinations, springs may be arranged in TRANSFER COMBINATIONS as shown in Figure 5. A transfer combination transfers a current or voltage from one circuit to another. The switch shown in Figure 5 uses two armature springs. Spring number 2 is an armature spring which is actuated indirectly through a buffer rod. Note that a voltage or current may be transferred from the circuit through springs 2 and 3 to a different circuit through springs 2 and 1. Thus, while two springs may be used for make and break combinations, at least three are required for a transfer combination. The transfer combination between springs 1, 2 and 3 in Figure 5 is a BREAK-BEFORE-MAKE COMBINATION. This means that, as the armature spring number 2 moves upward, the circuit between springs 2 and 3 is broken, or opened, before the new circuit between springs 1 and 2 is made, or closed. Another type transfer combination exists between springs 4, 5 and 6. This arrangement is called a MAKE-BEFORE-BREAK COMBINATION. As armature spring number 6 moves upward, the circuit between springs 4 and 6 is closed. Note that at this point all the contacts on springs 4, 5 and 6 are in the closed position. The overtravel, however, breaks the circuit between springs 4 and 5, and the new circuit is provided between springs 4 and 6. This operation is also called a CONTINUITY TRANSFER since voltages and currents may be transferred between two circuits without a momentary interruption which is the usual characteristic of transfer switches.



Snap-acting Switches

Another type of switch, and one which is in more general use, is the SNAP-ACTING SWITCH. There are a great many types of snap-acting switches. Among the common examples shown in Figure 6 is the TOGGLE switch used on electronic instruments. The toggle switch mechanism is often used with a variety of actuator types. In addition to the familiar bat-handle actuator, the ROCKER-TYPE actuator is used where a more flush-mounting switch is needed. The actuator assembly of snap-acting switches is more complex than the simple lever used in leaf switches. Basically, all snap acting switches use closely spaced, flat contacts which are snapped together or apart by a spring mechanism. A typical snap-acting mechanism is shown in Figure 7.

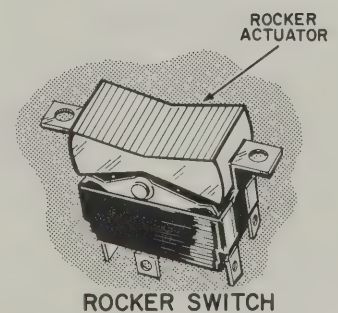
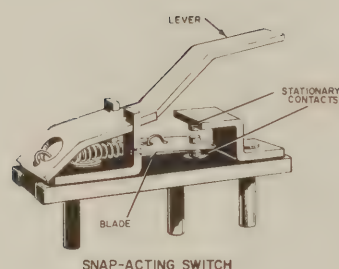


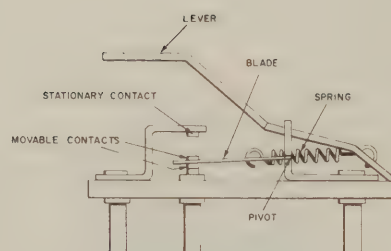
Figure
6

Snap-acting switches most often perform simple on-off and transfer switching. Complex switching is not often used, since there are a variety of other types of switches for these purposes. Thus, snap-acting switches are often SPST, SPDT, and DPDT mechanisms. A SPDT (single-pole, double-throw) switch has ONE electrical circuit (and therefore, one pole) and TWO closed switch positions (throws). The actuator may be positioned in either one of the two positions, and thus the switch is a SPDT switch. Figure 7 shows two views of such a switch. With the lever as shown, the lower movable contact rests against the lower stationary contact. Pressing downward on the lever

expands the spring, causing it to snap upward. The upward motion of the blade forces the upper movable contact against the upper stationary contact. When pressure on the lever is removed, the switch blade and contacts return to the position shown in Figure 7.

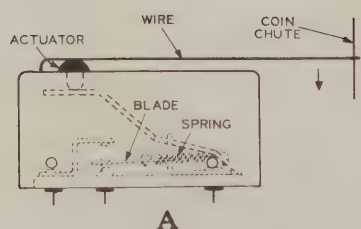


A



B

Figure 7



A

Figure 8

A DPDT (double-pole, double-throw) switch has two electrical circuits (poles) with contacts which are GANGED. That is, the switch actuator moves both poles simultaneously. Each of the poles has two closed positions (throws), and the switch is a DPDT switch. One, two and three positions are common, with two extreme (upper and lower) and one center position available. The center position is often neutral, thus providing an open circuit until the switch is operated. These switches, typical of toggle mechanisms, may be spring loaded to remain in the last position until the handle is again moved or it can be designed to return to the center or one of the other positions when the operator releases the handle. A useful feature of snap-acting toggle and rocker switches is the visual indication of position. Operators of electronic equipment are able to notice at a glance the on or off condition of rows of these switches.

Toggle and rocker switches are designed for manual operation by human hands. Another type of switch uses a snap-action mechanism which is often used with machines and automatic equipment. These switches use a PRECISION SNAP-ACTING mechanism which is similar to the mechanism used in toggle switches. The basic difference is that the actuator assembly has predetermined and accurately controlled travel distances. A precision snap-action switch requires contact separations which normally are no more than $\frac{1}{8}$ inch. The spring mechanism uses a coiled spring similar to the one of Figure 7.

Precision snap-acting switches use a variety of actuator lever devices as required for automatic machine operation. Some of these are shown in Figure 8. The same basic switch is shown as a coin switch in Figure 8A, a cam switch in Figure 8B, a slide switch in Figure 8C, and a plunger switch in Figure 8D. Coin switches are used in automatic vending machines and telephones. A long wire arm is used which protrudes through a coin chute, and is easily deflected by dropping coins. A lightweight, easily moved lever is used on the coin switch. A pushbutton actuator is used on each of these switches. The cam switch responds to rotary actuator motion. Thus, a rotating machinery motion actuates the switch each time the high part of the cam passes over the switch lever arm. The slide switch responds to linear motion through wedge action. The actuator lever arm is equipped with a roller. Slide switches such as this are used to sense items being transported along a conveyor. This type of switch is also used as a limit switch to sense the limit of travel in moving machinery parts.

Precision snap-acting switches can be made small and sensitive to light pressures. For this reason, many of these specialized switches are called **MICRO-**

SWITCHES. Often, the applications for microswitches are also specialized. Microswitches can be placed within the moving machinery of industrial equipment to indicate the position of various parts. Microswitches are also used to control various internal functions of complex machinery. Since the switch can be placed inside the equipment with very little space requirement, a moving part that touches the switch can be identified and controlled. The control function may be to signal the start of another machine cycle, or to initiate a new sequence. In this way, microswitches are useful in providing information concerning the operation of complex equipment.

Rotary Wafer Switches

Rotary wafer switches, or simply wafer switches, are another type of electro-mechanical device used to control electric circuits. Wafer switches are, however, quite different from those previously discussed. The most obvious difference is the type of actuation required. Wafer switches require some type of rotational (angular) force in order to make and break the contact mechanism. Because of this, wafer switches are most often manually actuated. They may, however, be actuated in other ways, such as by motors or solenoids.

Wafer switches are often front-panel mounted on electronic equipment where operators can easily reach them. The wafer switch is valuable in electronic applications because of its complex switching capabilities. Figure 9A shows a small wafer switch. The main parts of the switch are the shaft, frame, detent assembly, rotor blades, and contacts. The frame consists of the through bolts, spacers and wafers. The shaft, frame, and detent assembly make up the actuator mechanism. The rotor blades and contacts comprise the contact assembly.

The frame assembly is essentially expandable. That is, the switch may have any number of wafers by using longer through bolts and additional spacers. The rotor blades and contacts attached to the wafers are ganged to the common shaft. Thus, when the shaft is turned, the rotor blades on all the wafers rotate together.

Figure 9B shows a single deck, including the wafer, rotor blade and contacts. In the example shown, the rotor blade has on it a single movable contact. Thus, this deck, if used in a complete switch, would be a **SINGLE POLE** switch. The wafer also has attached to it 12 fixed contacts. The movable contact can be rotated to make contact with any one of the 12 stationary contacts. This switch is therefore a 12 position switch. By adding additional contacts (movable or fixed) and additional wafers, a great number of positions can be incorporated in a single switch. This is one of the advantages of a rotary wafer switch.

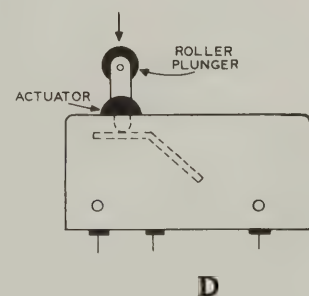
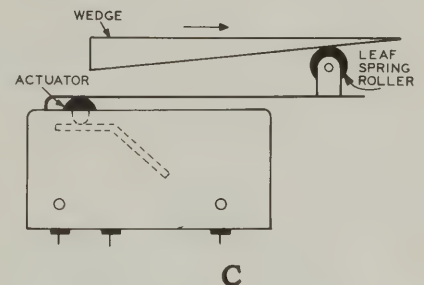
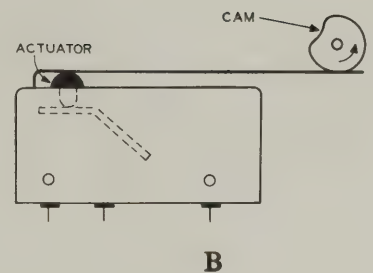


Figure
8

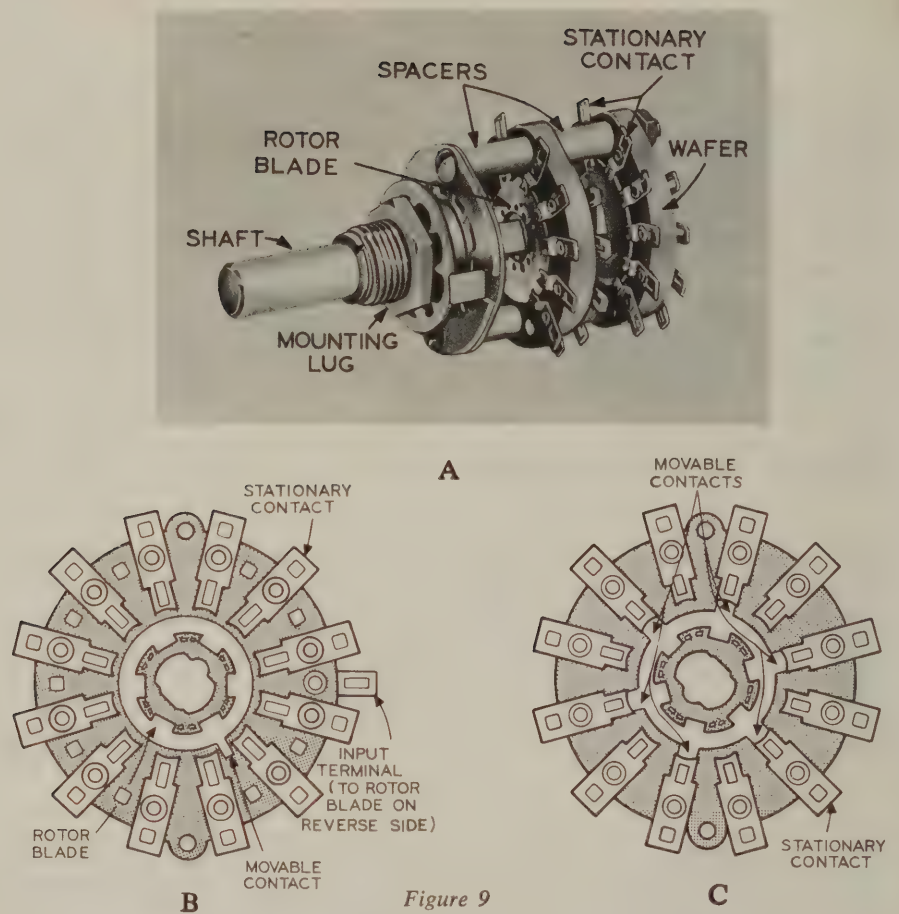


Figure 9

Figure 9C shows another single deck. This one is similar to the previous switch, but the contacts are arranged differently. Note that this time there are 6 movable contacts attached to the rotor blade. There are also 12 stationary contacts. In actual use, the wafer in Figure 9C would probably be used in a 2 POSITION switch. Since the fixed and movable contacts alternate around the wafer, a common application would be to switch 6 circuits simultaneously.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

SWITCHES

LEAF SWITCHES

SNAP-ACTING SWITCHES

ROTARY WAFER SWITCHES

1. Simple SPST switches, as shown in Figure 2, often used in power circuits that complete or interrupt a circuit are called _____.
2. Of the two main parts of any switch, the part which makes, breaks, or transfers a circuit is called the _____ assembly.
3. The spring in a leaf switch operated by the actuator assembly is an _____ spring.
4. The contacts on springs 2 and 3 of Figure 5 are in the (a) N.O., (b) N.C., (c) transfer position.
5. The minimum number of springs required for a transfer combination in a leaf switch is two. True or False?
6. Snap-acting switches are used mostly for (a) on-off switching, (b) telephone circuits.
7. Precision snap-acting switches are widely used in industry for the control of _____.
8. If used in an actual switch, and all contacts were utilized, how many POSITIONS might the switch in Figure 9B have?

RELAYS

RELAYS are electromechanical devices that use electromagnetic coils to perform remote switching operations. Figure 10 illustrates the basic relay principle. An electromagnetic coil is mounted inside a U-shaped housing which has one hinged arm. The hinged member is called the **ARMATURE**. The coil is mounted firmly to the housing with a mounting screw through the core at the end opposite the armature. This mounting member is the **HEELPIECE**. Actually, since the housing is a U-shaped member with only two separate pieces, the entire stationary portion which provides support for the coil may be thought of as the heelpiece. The entire assembly of Figure 10 provides the basic components of an **ARMATURE RELAY**. As shown in Figure 10, the housing provides a magnetic circuit for the electromagnet from pole to pole of the electromagnet core. The armature is normally held up, or away, from the electromagnet by spring action. The type of spring used may vary. Spring action is sometimes provided by the hinge construction. The coil and housing shown in Figure 10 are in the **DE-ENERGIZED** position. The armature is positioned away from the iron core by the spring pressure. The heelpiece, however, is securely fastened to the core at the other end. Thus, the housing provides a nearly complete **MAGNETIC CIRCUIT** for the magnetic lines of force when current is applied to the coil.

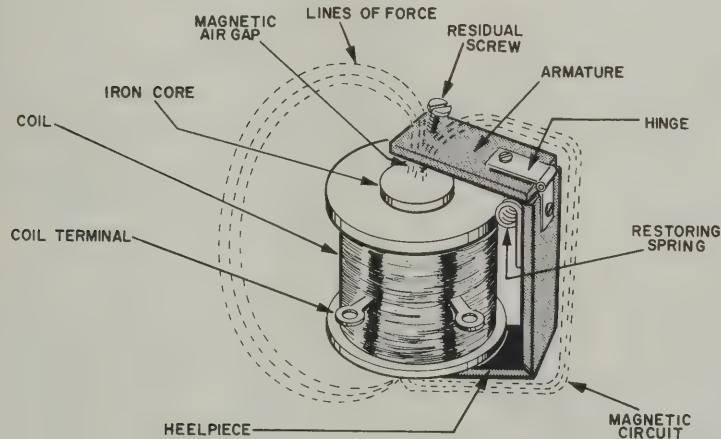


Figure 10

That is, the magnetic lines of force produced by the current flowing in the coil travel from one end of the coil, through the core and return externally through the housing and an air gap. The separation between the armature and the iron core is called the **AIR GAP**. The magnetic air gap is defined as **THE NONMETALLIC PORTION OF A MAGNETIC CIRCUIT**. When the coil is energized by applying current to the coil terminals, an electromagnetic field builds up around the coils as shown by the lines of force. The metal housing concentrates the magnetic lines of force. Most of the field produced by the electromagnet travels through the metal housing. As the

magnetic field builds up in the air gap, magnetic attraction overcomes the spring tension on the armature. When this happens, the armature is pulled toward the iron core, and the relay is ENERGIZED.

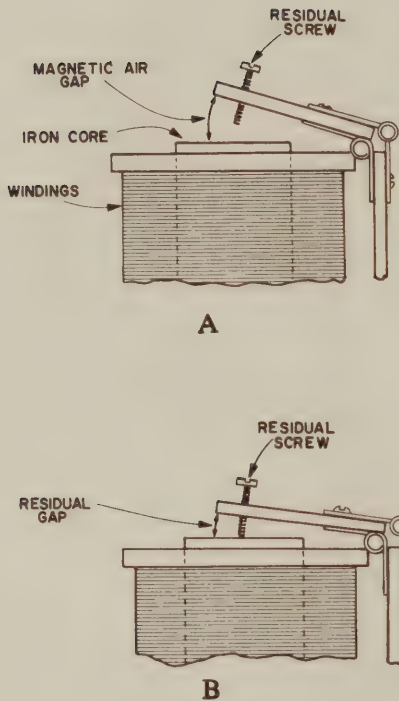


Figure
11

Generally, the armature is not allowed to completely contact the flat area of the exposed core. The reason for this is because the core of an electro-magnet becomes partially magnetized from repeated energization. This small amount of permanent magnetism is called RESIDUAL MAGNETISM, and would be sufficient to hold the armature against the core even after the coil is deenergized. Since the armature is supposed to return to the raised position due to spring tension when the coil is deenergized, residual magnetism causes an undesirable condition called STICKING. The effects of residual magnetism are reduced by preventing the armature from touching the core. This can be done in any of several ways. One of these is by using a RESIDUAL SCREW as shown in Figure 11A. Figure 11B shows the basic armature relay in the energized position. Note that the residual screw prevents the armature from completely contacting the iron core. Thus, a small magnetic air gap remains after the relay is energized. This is called the RESIDUAL GAP, and is defined as the gap between the core center and the nearest point on the armature WHEN THE RELAY IS ENERGIZED.

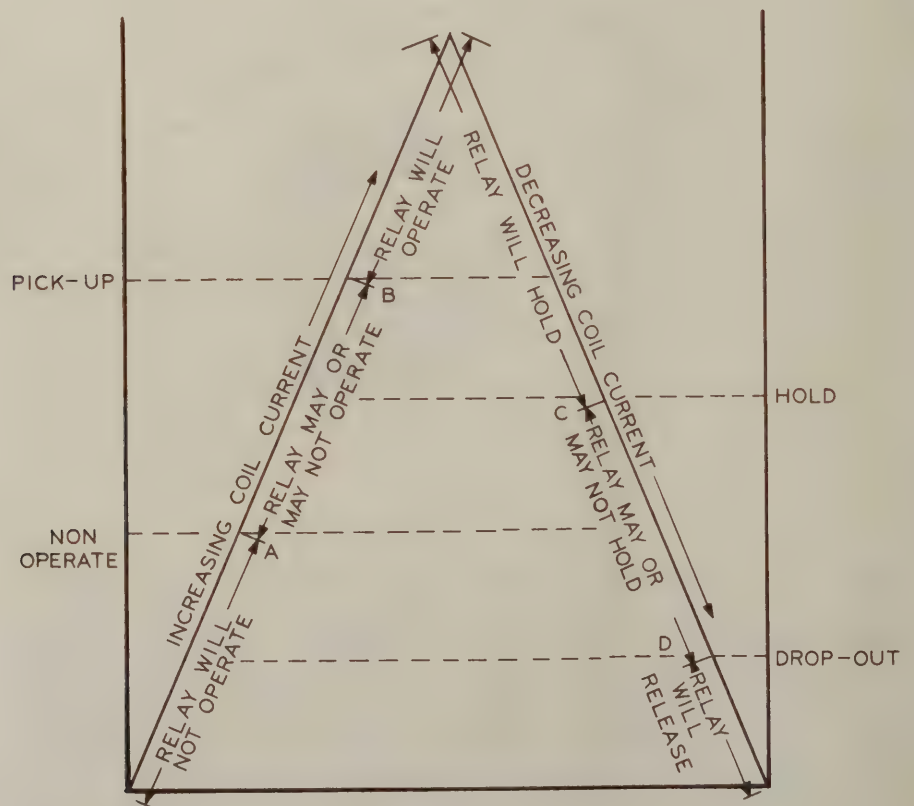


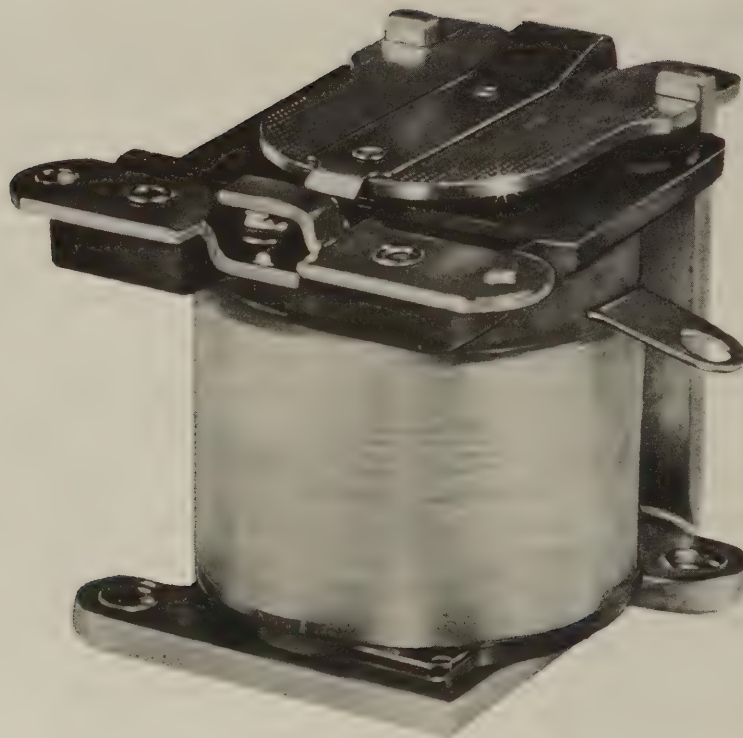
Figure 12

The magnetic air gap plays another important function in the operation of a relay. This is because an air gap, however small, greatly reduces the electromagnetic flux in a magnetic circuit. Since a relay armature is hinged, the magnetic circuit represents a VARIABLE RELUCTANCE to the electromagnetic lines of force. In the deenergized position, the magnetic air gap is large, and a large electromagnetic flux is required to overcome the restoring spring force. As shown in Figure 12, small current values, from the lower left corner of the graph to point A are insufficient to overcome these opposing forces, and the relay is represented as being in the NON OPERATE region.

As coil current is increased, a point is reached where the flux begins to overcome the reluctance of the magnetic circuit. The flux may be sufficient at this point to move the armature slightly, but without any definite motion toward the coil. The armature may, instead, chatter in response to slight current variations or transients through the coil. It is further possible that a transient current through the coil will momentarily increase the flux density, and cause the armature to move toward the coil. This is shown in Figure 12 in the "may or may not operate" region of the diagram between current values A and B. Note that the relay is still represented as being in the NON OPERATE area, since it cannot be definitely stated that the relay will energize with these values of current.

By increasing the current still further (above current level B) the flux density is built up sufficiently to overcome the reluctance of the magnetic circuit and the force of the restoring spring, and the relay will operate. The value of current, voltage or power at which a relay will ALWAYS operate is called the **PICK-UP VALUE**.

Once the armature has been operated, as shown in Figure 11B, the residual gap is quite small compared to the original magnetic air gap in Figure 11A. If, at this point it is desired to release the armature, again following the graph in Figure 12, it is shown that the relay will not release when the current is lowered to the pick-up value at current level B. Instead, the relay will definitely HOLD in the operate condition until level C is reached, which is considerably BELOW the pick-up value. This is because the residual gap is smaller than the original magnetic air gap, and the reluctance of the magnetic circuit has been greatly reduced. Therefore, the electromagnetic flux in the magnetic circuit is larger, and the armature remains down or operated. By decreasing the coil current further, an area between coil current levels C and D is reached which is similar to the area between levels A and B. Again, circuit transients and mechanical irregularities may or may not cause the armature to hold, but then again it might not hold. Only at a much lower current, at level D, can it be stated that the armature will always release. This point is labeled **DROP-OUT**, and is the coil current voltage or power value for which the armature will ALWAYS release. The relay is then deenergized.



Often called a dc plate relay, this relay uses the most basic construction and is suitable for most electronic applications.

Courtesy Potter and Brumfield, Div. AMF, Inc.

Thus it is seen that, for any relay, more current is required to energize a relay initially than is required to keep it energized. This, again, is due to the difference in the length of the magnetic circuit air gap between a deenergized relay and an energized relay.

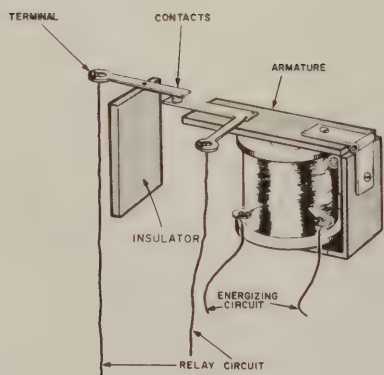


Figure 13

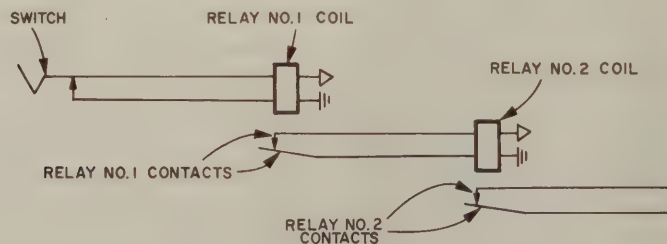


Figure 14

Figure 13 illustrates further the component arrangement of Figure 11A and 11B with the addition of contacts and circuit wiring. With these additions, the device becomes an armature relay. There are two separate circuits associated with any relay: the ENERGIZING or coil CIRCUIT; and the RELAY CONTACT, or trip CIRCUIT.

The energizing circuit operates the contact assembly, and in this way is similar to the actuator in a switch assembly. The relay contact circuit makes, breaks, or transfers one or more different circuits, and performs the same function as the contact assembly in an electromechanical switch. Relays often perform functions more complex than those of manual switches. For example, the energizing circuit of one relay may be controlled by the relay contact circuit of another relay, as shown in Figure 14. From interconnections such as this, a generalized definition for a relay is often used: A relay is "an electromechanical device which is operated by variations in the conditions of one electric circuit to affect the operation of other devices in the same or other circuits by the opening and closing of contacts".

General Purpose Relays

Relays are classed as to the functions they perform. Although many relays look alike, there are often very important differences. Such things as coil power requirements, contact current capabilities, transfer combinations and many other characteristics are carefully considered by relay users. In general, any relay which is not highly specialized and is available as a "stock item" by a manufacturer is considered to be a GENERAL PURPOSE RELAY. General purpose relays are sub-classified by type, such as: ac or dc, light duty or heavy duty, standard size or miniature, sealed or open.

Figure 15 shows two views of a light duty, general purpose relay. This relay, while considered light duty is somewhat ruggedized. It will withstand a degree of shock and vibration since the contact springs are relatively short and the contacts are large. Many types of relays do not work well when subjected to vibration. Relay contact assemblies must use flexible arms and moving mechanisms to effectively switch circuits. Vibration may cause the armature arms to bounce between the contacts and produce poor performance. Thus, the contact assembly characteristics are carefully selected to suit the application.

Figure 16 shows two views of another general purpose relay. This one would not be suitable for use in high vibration environments. The relay uses long leaf springs, and resembles the leaf switch discussed earlier. Vibration and shock would cause the springs to bounce, and the relay would produce intermittent circuit operation. A desirable feature of this relay is its low cost.

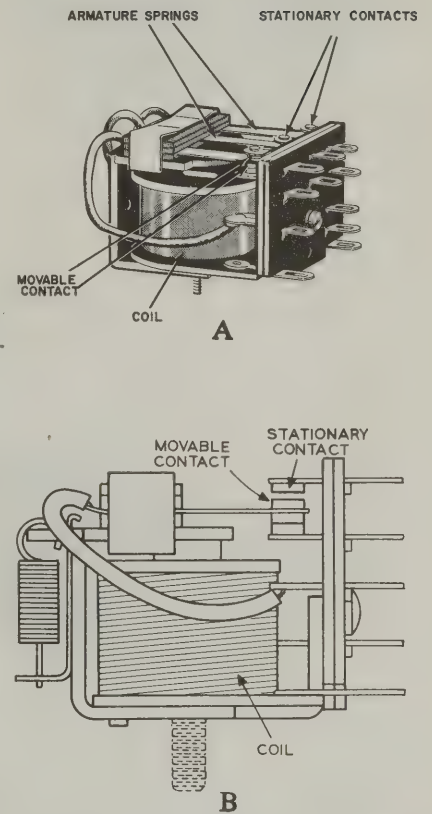


Figure 15

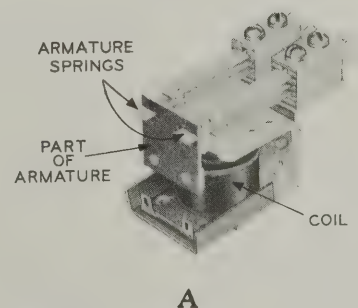


Figure 16

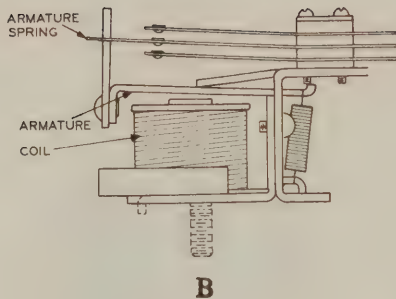


Figure 16

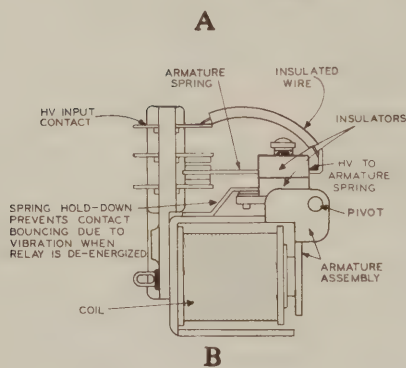
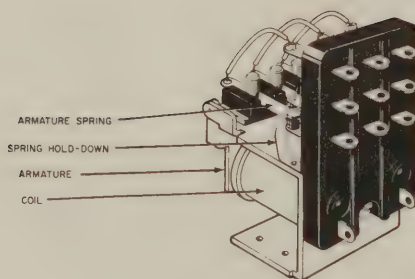


Figure 17

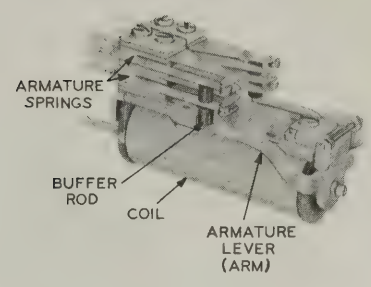
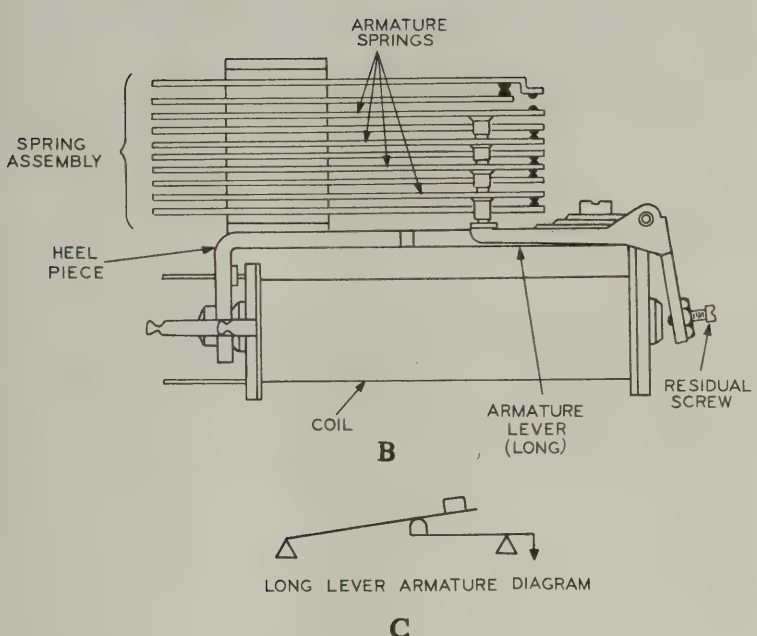
Figure 17 shows a third general purpose relay. This relay is considered heavy duty, and would operate under moderate amounts of shock. Since the relay is designed for heavy duty, the springs are thicker and heavier than most relays of this type. The springs are designed to carry relatively high currents, and the coil also will withstand prolonged or continuous energization.

This type of relay is sometimes classed separately from other general purpose relays and called a **POWER RELAY**. Power relays are actually extremely rugged general purpose relays. There are also various degrees of ruggedness required in power applications, and therefore some variations are required in the design of power relays. Basically, however, power relays must be able to switch circuits carrying ac "line" currents, and dc voltages from large dc power supplies. They are used in conjunction with high voltage transformers, power generating equipment, trucks, and other heavy duty equipment. Since power relays are used in the electrical systems of machinery, these devices must be able to withstand shock, vibration, and prolonged operation. Note also in Figure 17 that a hold-down spring is used which prevents the contacts from bouncing, or "chattering", when the relay is de-energized. This spring bends upward when the relay is energized, thus permitting normal operation when sufficient current is supplied to the coil to energize the relay.

Telephone-type Relays

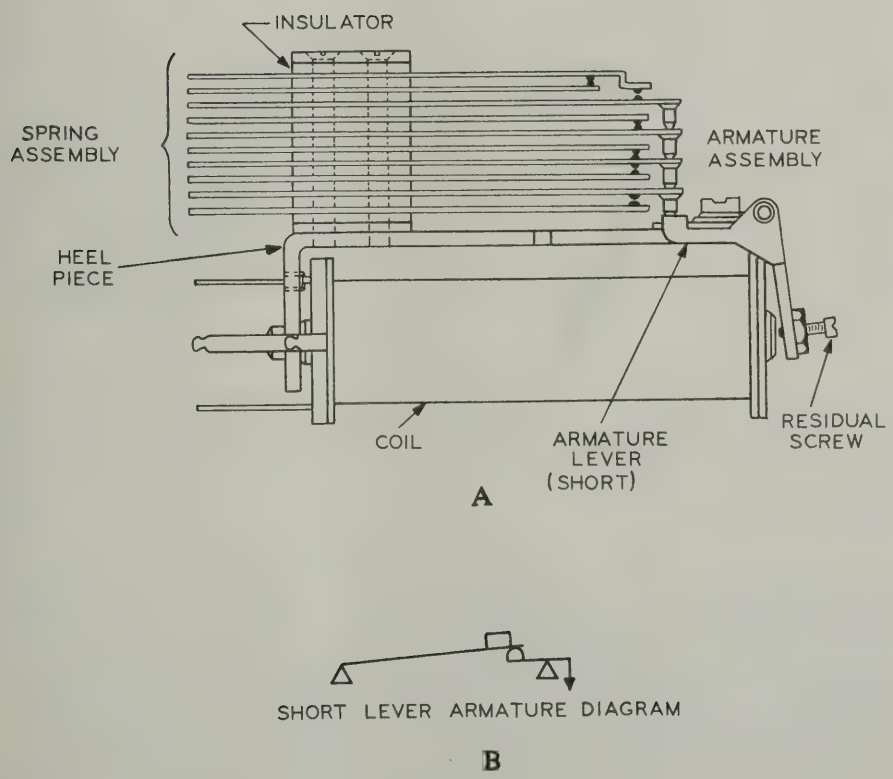
Relays are used extensively in telephone systems, and several types of relays have been designed specifically for the needs of telephone switching. One of the unique characteristics of a telephone relay is that the switched circuits must be VOICE GRADE. That is, the circuits carry audio signals rather than power or control signals. Since noise must be kept to a minimum, some very special considerations are given to the design of telephone relays. Although several different types are used in telephone systems, one type, shown in Figure 18, has been in use since the early days of telephony and is similar to the leaf switch discussed earlier. The distinguishing features of a telephone-type relay are often a heavy spring pile-up and a long coil. The spring pile-up may vary from the double set, as shown, to single sets. The shape of the springs also may vary.

Some relays use angled fixed contact springs, as shown in Figure 18A, and thus all the ARMATURE springs are in line, one directly above the other. The buffer rods, in the arrangement, are also in line, and the armature springs are able to move together. As shown in Figure 18A, this permits the use of a LONG ARMATURE and thus permits a shorter magnetic air gap. This means that the leverage is reduced, and the pileups used with long armature relays must be small. The advantage gained, however, is reduced current drain, since the air gap is smaller and the relay will operate with lower energizing currents.



A

Figure 18



B

Figure 19

Figure 19 illustrates a **SHORT ARMATURE** relay. Note that the pile-up is quite heavy. That is, there are a considerable number of springs, all of which oppose the armature with spring tension when the relay is operated. In order to operate this relay, the buffer rods must be placed at the ends of the armature springs to obtain the most leverage. The armature arm must be shorter, and the travel distance greater. In order to obtain the required travel distance, the air gap must be larger. This is a disadvantage, since greater energizing current is required. However, this is unavoidable, since the relay must operate a heavy spring pile-up.

Another characteristic of telephone relays is the long coil. These coils have a large number of turns, and are magnetically very efficient. Since telephone systems use large numbers of relays, the special design of telephone relays centers around the reduction of power consumption and the economical use of materials. Other special operating characteristics of telephone relays rely upon the unique shape of the coil. Among these interesting characteristics are time delays.

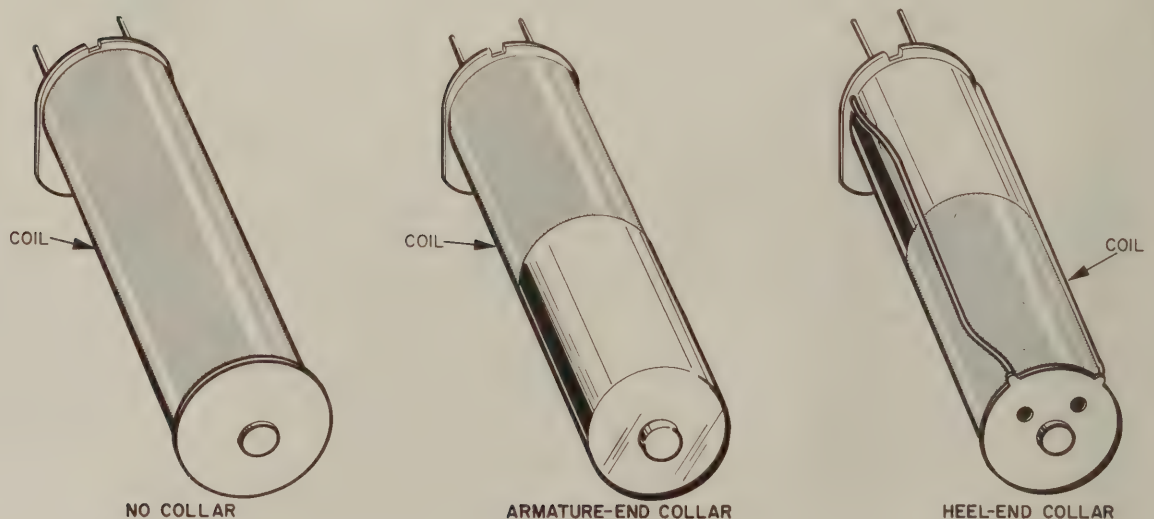


Figure 20

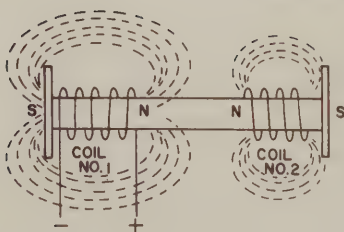
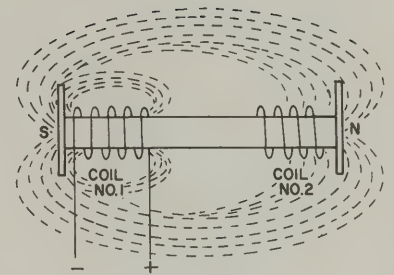


Figure 21

Figure 20 shows three telephone relay coils, two of which have a solid copper **COLLAR** attached. A collar attached as shown to either the armature or heel end affects the pickup and dropout time of a relay. That is, a relay will be either SLOW OPERATING (S. O.), or SLOW RELEASING (S. R.) depending upon where the collar is placed. Figure 21 illustrates the flux paths for relay coils with collars attached. The armature end collar in Figure 21A causes the relay to be slow operating (S. O.). To understand the operation of a relay with a copper collar at one end, think momentarily of the

collar as a standard coil of wire. There would then be two coils placed side by side on the same core as shown in Figure 21A. When current is applied to coil 1, an electromagnetic field builds up around it and cuts across the windings of coil 2. This produces a counter electromotive force by induction, in coil 2. The electromagnetic field now building up around coil 2 opposes the field around coil 1.

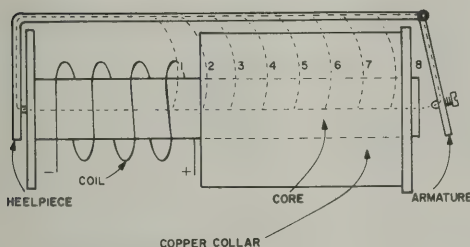
The wire in a coil has resistance. This resistance limits the CEMF that is developed in coil 2 by induction from coil 1, and the electromagnetic field around coil 2 is small, as shown. Figure 21B shows the condition of the two coils a few moments after coil 1 has been energized. After the coil 1 flux has stabilized and is no longer expanding, the CEMF in coil 2 stops and the field around coil 2 disappears. The electromagnetic field around coil 1 is then unopposed, and uses the entire length of the core for its external magnetic circuit.



B

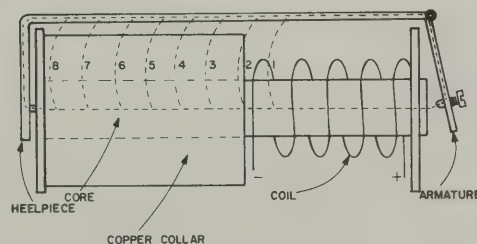
Figure 21

Figure 22 shows the same basic two coil configurations with the addition of a relay housing, and the substitution of coil 2 with a solid copper collar. The collar differs from coil 2 in that it is highly conductive. The collar, with its much lower resistance, acts as if all the turns of wire in a coil were shorted together. Therefore, the collar generates a large electromagnetic field when the flux from coil 1 cuts across it. In fact, the field is as large as the field around coil 1. Note that in Figure 22A, the electromagnetic field PROGRESSES through the collar, and the field around the collar does not appear all at once. Since the collar acts as a coil with a very large number of turns, it takes time for the coil field to travel through the collar. This is indicated in the Figure by the lines numbered 1 through 8, which represent the expanding lines of force. At each point (number 2 through 8) the expanding lines of force generate a CEMF in the collar, and cannot progress until the resulting field in that portion of the collar collapses. This happens all along the length of the collar until point 8 is reached. In Figure 22A, point 8 is at the arma-



- 1 MAGNETIC FLUX PATH IMMEDIATELY AFTER COIL-CIRCUIT CLOSURE
- 2 MAGNETIC FLUX PATH A MOMENT LATER
- 8 FINAL FLUX REACHES AIR GAP TO CAUSE OPERATION

A



- 1 MAGNETIC FLUX PATH IMMEDIATELY AFTER COIL-CIRCUIT CLOSURE
- 2 MAGNETIC FLUX PATH A MOMENT LATER
- 8 FINAL FLUX LINKS COLLAR, COLLAR NOW CAN DELAY RELEASE

B

Figure 22

ture. The armature, remember, is part of the external magnetic circuit. When the coil flux reaches this point, the magnetic circuit is complete, with the exception of the relay air gap. With the flux concentrated in the external magnetic circuit at point 8, there is sufficient attraction between the core and the armature to overcome the air gap and pull the armature down. The time required for all of this to happen is considerably longer than would be required if there were no collar on the relay. A large collar, placed near the armature of a relay, can delay relay operation for several seconds. The length of the time delay is determined by the size of the collar, and by other factors such as the coil size, air gap distance, and magnetic circuit characteristics.

When current is removed from the coil, an electromagnetic field again builds up around the collar. Since the collar is placed near the armature, the field maintains flux linkage in the residual gap, and the armature remains operated. Thus, AN ARMATURE-END COLLAR CAUSES A RELAY TO BE BOTH SLOW OPERATING AND SLOW RELEASING.

Figure 22B shows the effect of a heel-end collar on a relay. This time, note that the collar is not close to the armature, but the coil is. When the coil is energized, the electromagnetic flux builds up rapidly, and uses PART of the relay housing as the external magnetic circuit. The armature is included in this circuit, and the relay will operate with practically no delay. There is some small delay, however, since the magnetic circuit is not as good as it would be if the flux could travel through the heelpiece. As shown before, the flux cannot travel through the heelpiece initially since it must travel through the collar first. This delay is actually very small, and we may say that a heel-end collar on a relay DOES NOT AFFECT OPERATION (energizing) of the relay. When current is removed from the coil, the electromagnetic field builds up again around the collar. Since the relay is already operated, the armature touches one pole, and the field generated by the collar has an excellent external magnetic circuit. Again, flux is maintained in the residual gap by the

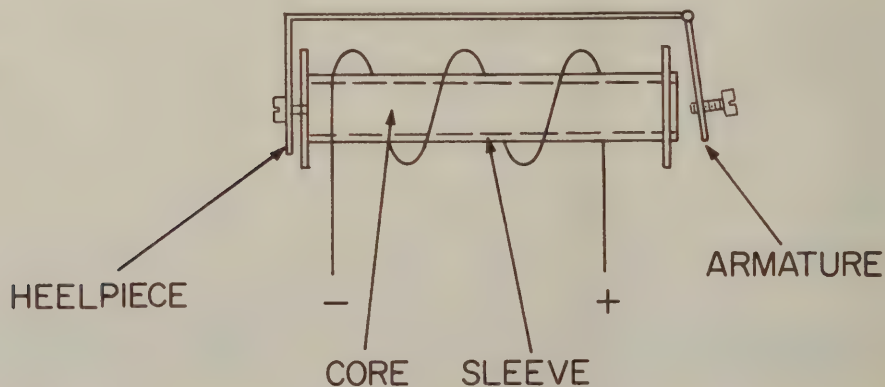


Figure 23

collapsing coil field, and the relay remains operated for a time after coil current has been removed. Thus, A HEEL-END COLLAR DOES NOT AFFECT SUBSTANTIALLY THE OPERATION (energizing) OF A RELAY, BUT CAUSES A RELAY TO BE SLOW RELEASING (de-energizing). Relays having a collar at one end are sometimes called SLUGGED RELAYS.

Collars are widely used in telephone switching systems to provide timing. The collars are solid copper and quite large. Figure 23 shows a method for reducing the amount of copper required to provide slow releasing. The core of the relay is completely enclosed with a copper SLEEVE, and the coil is wound along the entire relay length over the sleeve. This type of relay, called a SLEEVED RELAY, is similar to a heel-end collar relay in operation. That is, the operation (pick-up) of the relay is substantially unaffected. This is because the coil is near the armature and the heelpiece. The magnetic circuit is completed upon energizing the coil with little interference from the sleeve. When the coil is deenergized, however, the armature is already operated and is close to the relay core. Remember that a smaller magnetic field is required to hold a relay operated than to pull the armature down initially. As the coil field collapses, a field is generated around the sleeve THROUGH THE CLOSED MAGNETIC CIRCUIT which is strong enough to hold the armature down, thus keeping the relay energized. When the electromagnetic field around the coil has collapsed completely, the field around the sleeve will disappear and the armature will restore.

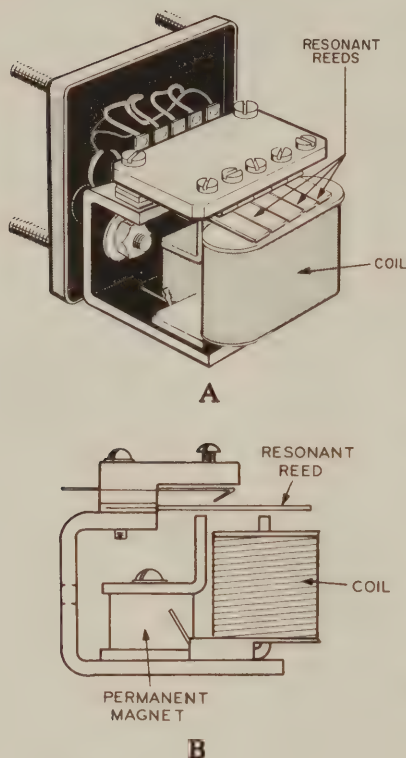


Figure
24

Resonant Reed Relays

A RESONANT REED RELAY is an electromechanical device that acts as a frequency-sensitive switch. This relay consists of a number of precisely-tuned metal reeds which are free to vibrate as single tines of a tuning fork, as shown in Figure 24A. The reeds are attached at one end to a permanent magnet (shown in Figure 24B), and thus become temporarily magnetized (biased). An electromagnetic coil is placed in close proximity to the free ends of the reeds. When an ac signal is applied to the coil, the electromagnetic field fluctuates in response to the current flowing through the coil. The frequency at which the electromagnetic flux varies determines which reed will vibrate. The reeds, being of different lengths and thicknesses, are sensitive to different frequencies. This specific frequency that each reed responds to is the RESONANT FREQUENCY of that reed. Thus, the varying electromagnetic field coincides with the mechanical resonant structure of only one of the reeds at any given time. When the reed vibrates with sufficient amplitude, the reed touches a switch contact placed between the reed and the electromagnet. By vibrating against the contact, a current flows through the reed and the switch contact once during each cycle of vibration. This produces an intermittent, or pulsed, dc output.

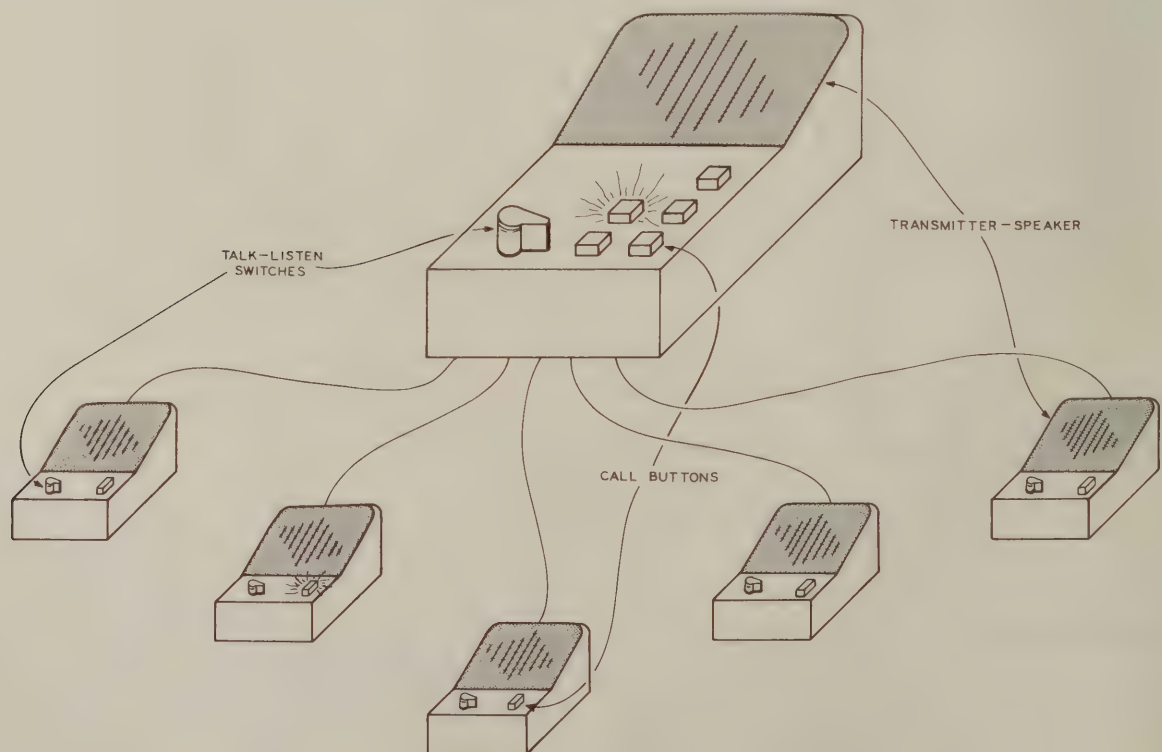


Figure 25

The following Practice Exercise questions cover the subjects which you have just studied. They are:

RELAYS

GENERAL PURPOSE RELAYS

TELEPHONE-TYPE RELAYS

9. The movable, hinged portion of a relay housing is called the (a) armature, (b) heelpiece.
10. The purpose of the residual screw is to prevent the condition called _____.
11. The value of current, voltage, or power at which a relay will always energize is the _____ value.
12. The value of current, voltage, or power at which a relay will always deenergize is the _____ value.
13. The circuit associated with any relay that operates the contacts is the (a) energizing circuit, (b) relay contact circuit.
14. General purpose relays which can withstand vibration are (a) ac, (b) heavy duty, (c) dc.
15. Relays which are often used in high voltage transformer circuits are called _____ relays.
16. An armature-end collar causes a relay to be (a) slow operating only, (b) slow releasing only, (c) both slow operating and slow releasing.
17. A heel-end collar does not substantially affect the OPERATION (energizing) of a relay. True or False?

Resonant reed relays are commonly used to "Ring" a desired intercom- receiver from a central intercom station while using audio circuits which are common to all the units. This is shown in Figure 25. The central operator depresses a call button which lights up and applies one of several preset frequencies of alternating current to the input lines of all the substations. Each substation has a resonant reed relay tuned to a different frequency. When the reed associated with this frequency vibrates, that substation is alerted. The relay output can be used to operate a buzzer and light a panel light in the desired remote intercom receiver. The remote intercom operator turns the Talk-Listen switch to the "listen" position, and is able to receive the message. A remote intercom unit operator can also initiate a call by depressing the call button on the remote unit. This lights up the call buttons on both the remote and central units. The remote operator then places the Talk-Listen switch in the "talk" position, and can then transmit a message. The system shown does not use call suppression, and therefore both the transmitting and receiving intercom call buttons will light up when the substation call button is depressed.

MAGNETIC REED RELAYS

A basic MAGNETIC REED RELAY consists of two flat, overlapping, low reluctance metal reeds separated by a small air gap and sealed in a glass envelope. This is shown in Figure 26. Used with suitable circuitry, the reeds will make, break or transfer an electric circuit under the influence of a magnetic field. Reed relays can be actuated by either PERMANENT MAGNET or ELECTROMAGNETIC fields. A magnetic field causes the reeds to move together by magnetic induction when the flux density in the air gap between the reeds overcomes the spring separation resistance of the reeds.

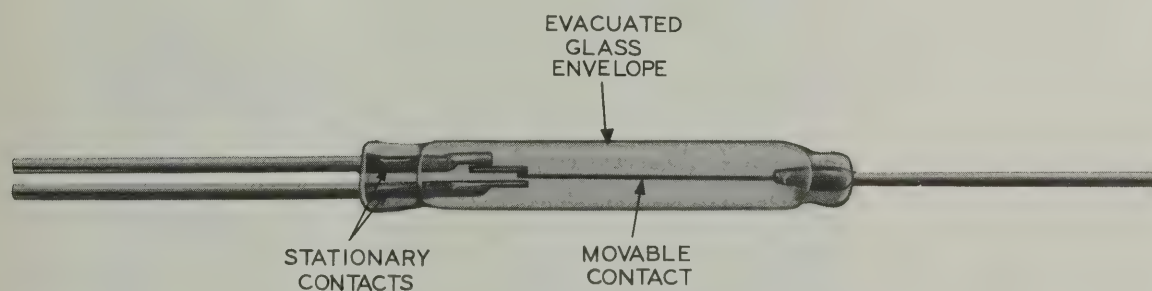


Figure 26

Reed relays may be actuated through several combinations of magnet and reed positions. These are shown in Figure 27A-F. Note that in all except Figures A and F, more than one pickup (PI) and dropout (DO) point occurs as the magnet is moved in relation to the reed. These multiple points are caused

by the combined variations in the magnetic field, which is a characteristic of the magnet, and variations in the magnetic circuit, which is a characteristic of the two reeds and the air gap between them. In the figures, the reed will close when the dashed lines near the PI-DO symbols are directly opposite the dashed line shown through the magnet.

Figure 27A illustrates permanent magnet operation using PERPENDICULAR ACTUATION. This arrangement has the magnet parallel to the switch and the magnet is moved perpendicular to the switch. It gives the most positive action using any given magnet. It also gives only one closure of the reed switch as the magnet is moved in a single direction toward the envelope.

Figure 27B illustrates PARALLEL ACTUATION. This method of actuation yields two closures when approaching the capsule from one of the ends. Note

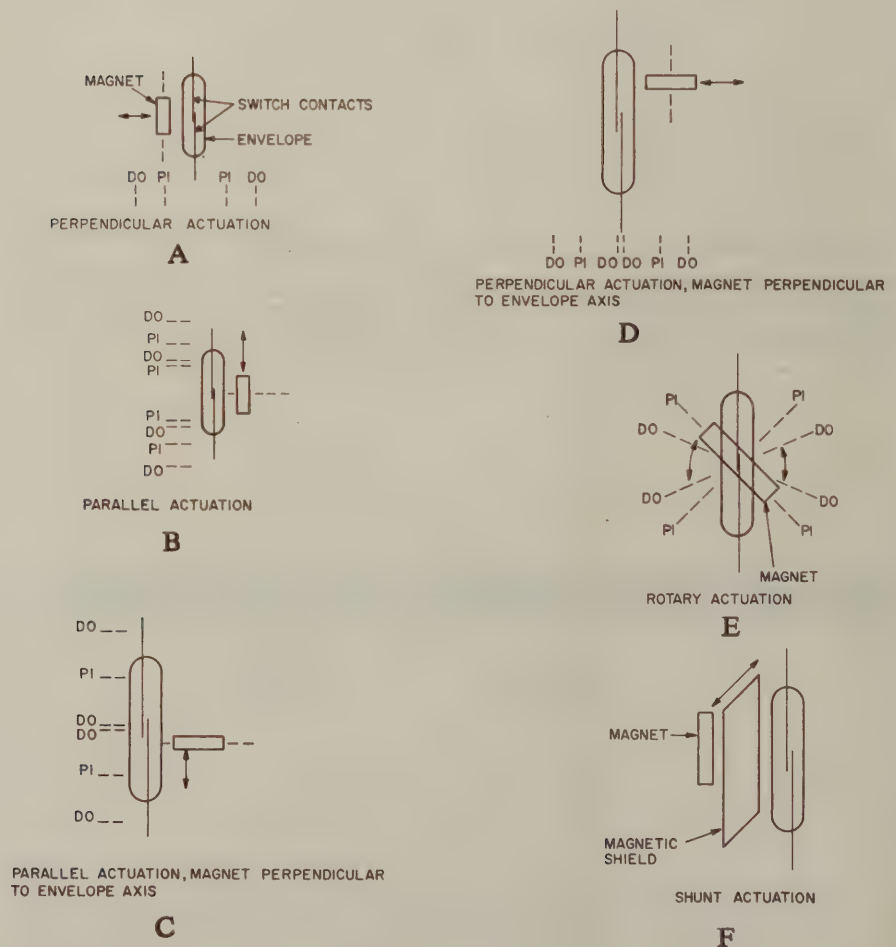
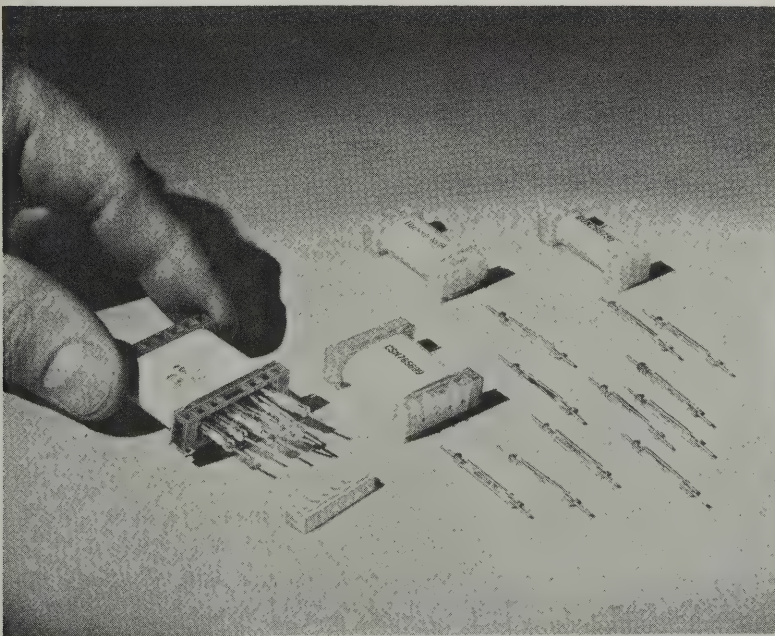


Figure 27

that actuation is the same when approaching the capsule from the other end. Parallel actuation is used where small travel distances are required between PI and DO.

Figure 27C shows parallel actuation when the magnet is perpendicular to the envelope axis. This arrangement results in one reed closure for each half of the envelope. The advantage of this type of operation is the close control of the dropout point. Note that dropout occurs when the magnet pole is directly in line with the air gap. This is because the reeds lie perpendicular to the magnetic field, and the air gap is not required for completion of the external magnetic circuit.

Figure 27D shows perpendicular actuation with the magnet perpendicular to the envelope axis. In this figure, visualize the magnet passing in front of or behind the envelope. This PI-DO pattern appears similar to the one for Figure 27C. However, the dropout point here is dependent upon the length of the magnet, and the magnet's distance away from the air gap. Note that in the method of D, dropout occurs when the magnet is in line with the air gap. This occurs for the same reason that was given for C. The air gap is not utilized in the external magnetic circuit when the magnet is in this position.



Magnetic reed relay capsules are often ganged together into multiple-pole modules, thus switching several circuits simultaneously when the coil is energized.

Courtesy IBM Corporation

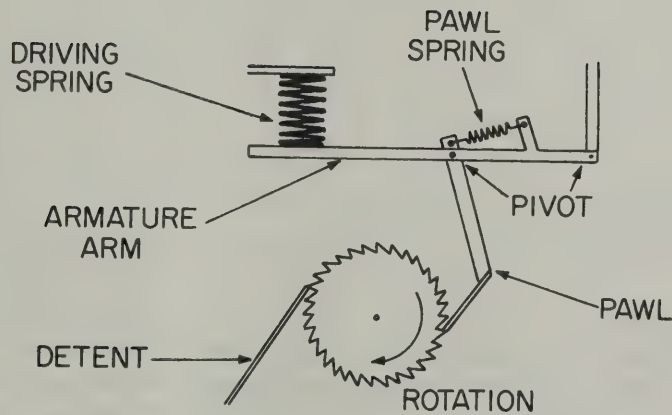
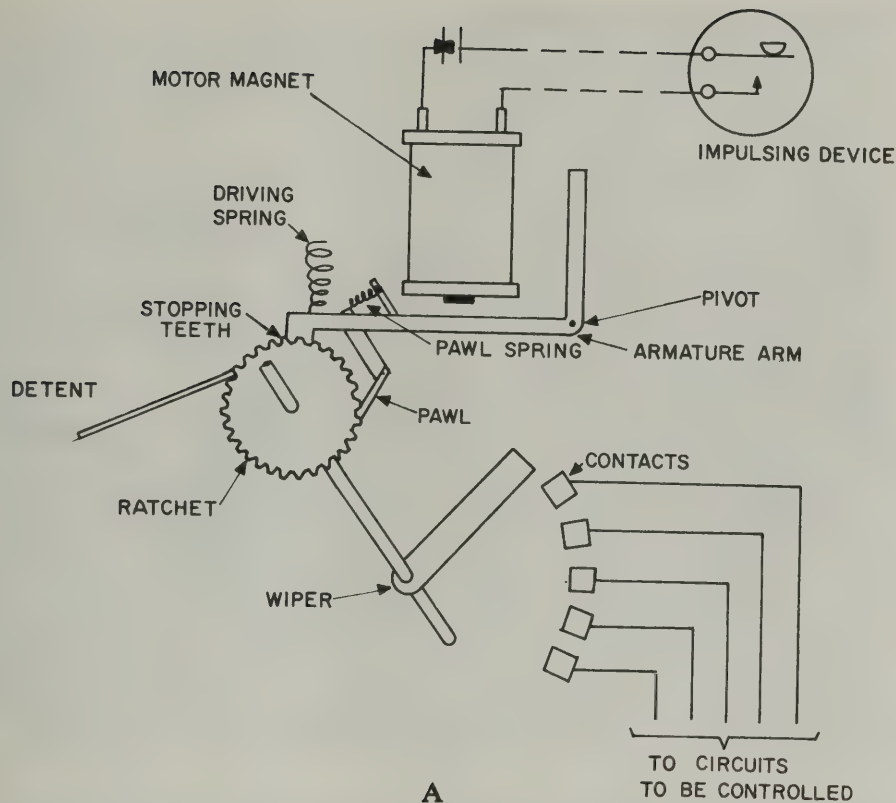
Figure 27E illustrates ROTARY ACTUATION. This arrangement results in reed closure twice during each revolution of the magnet. It can also be used as a positive on-off control when magnet movement is limited to 90° rotation.

Figure 27F illustrates shunt actuation. A reed relay can be controlled by the use of a magnetic shield inserted between the envelope and the magnet. This technique is used in applications where movement of the magnet relative to the reed envelope is not practical.

Reed relays can also be made with additional elements inside the envelope, thus permitting TRANSFER switching. A reed relay having a transfer combination is shown in Figure 26. In this relay, a single long reed functions as a movable armature. The armature reed maintains contact with one of the short reeds in the unoperated condition. This is a normally closed (N.C.) combination. When a permanent magnet is brought near the NORMALLY OPEN contact, the armature reed breaks the N.C. circuit, and makes (closes) the N.O. circuit.

STEPPING SWITCHES

An electromagnet can be used with some additional mechanical components to perform another type of switching. A simplified illustration of the mechanical linkage is shown in Figure 28. This device is a ROTARY STEPPING SWITCH and performs switching by moving a wiper arm across a set of contacts. The actuator assembly of a stepping switch is called a MOTOR MAGNET, and is similar to the electromagnet used as the actuator of a relay. Also similar to the operation of a relay is the use of an armature which is attracted by the motor magnet, thus operating the switch. A number of additional mechanical components are connected to the armature. The main component in the contact assembly is a RATCHET mechanism. Figure 28B is an enlarged section of the ratchet mechanism in Figure 28A. A ratchet is a toothed wheel which transmits a mechanical force in only one angular (rotational) direction. The element which applies the mechanical force to the wheel is the PAWL. The pawl, attached to the armature, is pivoted and held in the ratchet teeth under spring tension as shown in Figure 28B. In Figure 28A, the armature arm moves upward when the motor magnet is energized. Note, however, that during the upward motion, the armature compresses a DRIVING SPRING. When the motor magnet deenergizes, the energy stored in the compressed spring is released, and the armature is forced downward. As the armature moves up and down, the pawl applies downward force on the ratchet wheel. When the armature moves upward, the pawl does NOT apply driving force to the ratchet. Instead, the pawl merely rides over the ratchet teeth in preparation for the next downward motion. During the downward motion of the armature, the pawl pushes against the ratchet tooth, moving it.



B
Figure 28

The switch contact assembly shown in Figure 29 consists of two basic parts: the BANK ASSEMBLY or simply the BANK and the WIPER ASSEMBLY. The bank is stationary, and may have up to 25 contacts on each level. The switch shown has 25 contacts per level on 11 levels. This makes up the entire bank. All the bank contacts are electrically insulated from one another. The wiper arms, one per level, are attached to the ratchet wheel on a common shaft, and are movable.

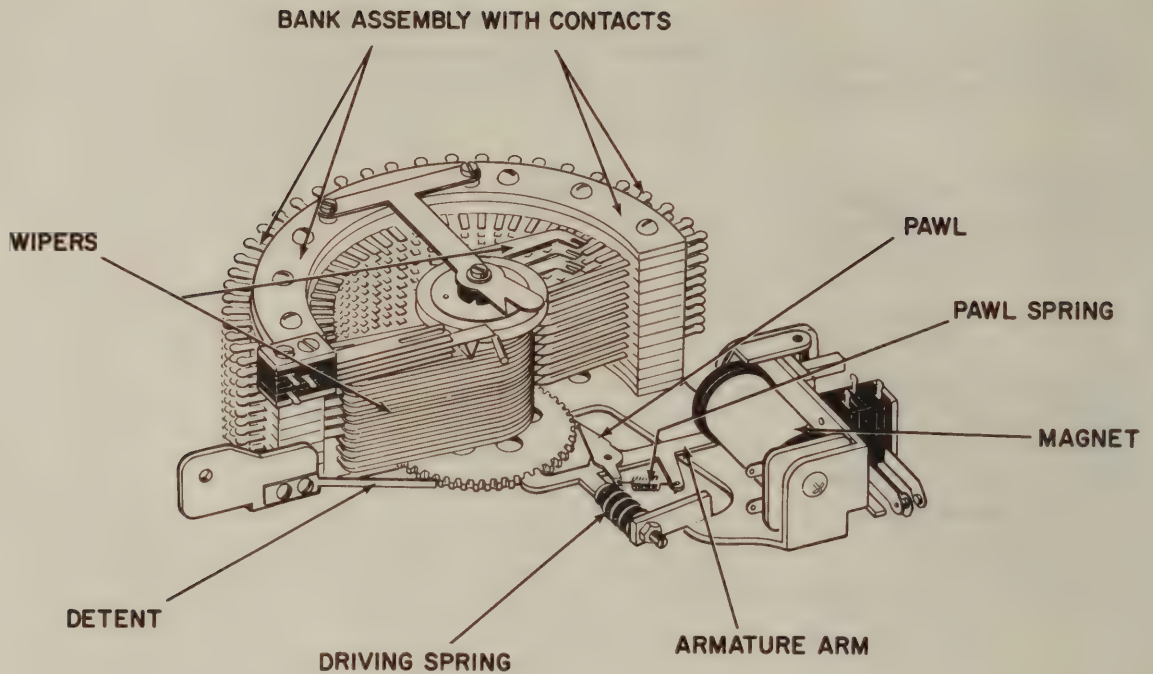
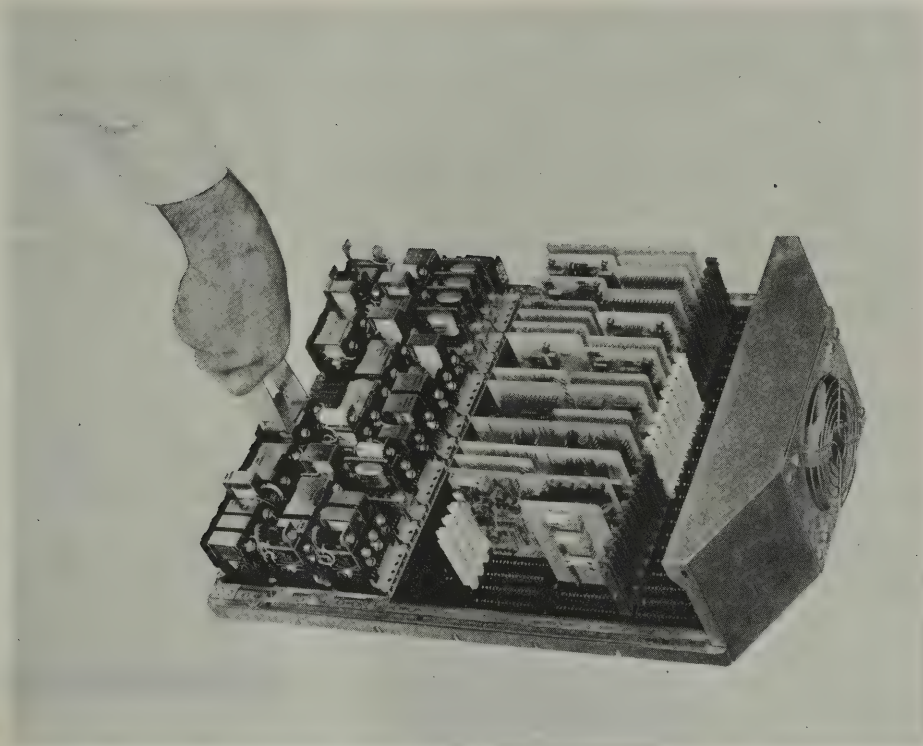


Figure 29

As the ratchet is turned, the wipers "step" from one bank contact to the next. Since the bank is semicircular, one set of arms moves across the band and just steps OFF the last row of contacts, while the other set of arms, opposed to the first by 180° , steps ONTO the first row of contacts. The wipers on different levels are electrically insulated from each other. However, since the 180° opposed wipers on any given level are electrically the same, the dual wiper system assures that the wipers are on the bank at all times. This also assures a shorter RESET cycle since the wipers require less than 180° travel to reach any previous bank contact.

Considering the number of contacts and wipers in a stepping switch, one can readily see that the contact friction is high. Considerable force is required to move the wipers across the bank. This large amount of force required to move the wipers across the bank is due to the MOVING FRICTION of the contact assembly. This means that the contacts present a large amount of friction which resists the actuation force while the wipers are in motion. As large as this resistance is, however, the friction which must be overcome to just START the wipers moving from a standing position is even greater. This resistance to actuator force is due to the STATIC FRICTION which resists the attempt to start the wipers moving. In other words, less force is required to KEEP the wipers moving than is required to START them moving. For this reason, most stepping switches use a driving spring to provide the force which moves the ratchet wheel and, ultimately, the wipers. Switches that employ the driving force of a spring are INDIRECT DRIVE switches. Springs



Another type of relay, called a wire contact relay, is widely used in communications equipment. This type of relay uses wires instead of leaf-type springs. This arrangement reduces the weight of the spring pileup and permits the relay to switch more sets of contacts. Here, a wire contact relay is removed using a relay puller.

Courtesy IBM Corporation

are used to obtain a **MECHANICAL ADVANTAGE** in operating the wipers of large stepping switches. Some switches use the driving force of the coil as it is energizing, and are **DIRECT DRIVE** switches.

SOLENOIDS

A **SOLENOID** is an electromechanical device which converts electrical energy into mechanical energy. Since solenoids are generally used for the express purpose of obtaining mechanical motion rather than to control another electric circuit, a solenoid is a **TRANSDUCER**. Note that the definition for a solenoid coincides with the definition for a transducer.

A solenoid consists basically of a coil of wire and a movable soft iron bar, called a **PLUNGER**. The plunger is free to move back-and-forth within the

length of the coil. A simple solenoid is shown in Figure 30. When the coil is energized, the portion of the plunger which is inside the coil greatly lowers the reluctance of the internal magnetic circuit for that portion of the coil. The other section of the coil, however, has an air core, and the reluctance of this part of the internal magnetic circuit is high. The electromagnetic field generated by current flowing through the coil is concentrated in the plunger. In order to reduce the high reluctance air portion of the magnetic circuit, the lines of force tend to draw the plunger further into the coil. This reduces the distance that the lines of force must travel through air, and further concentrates the electromagnetic field.

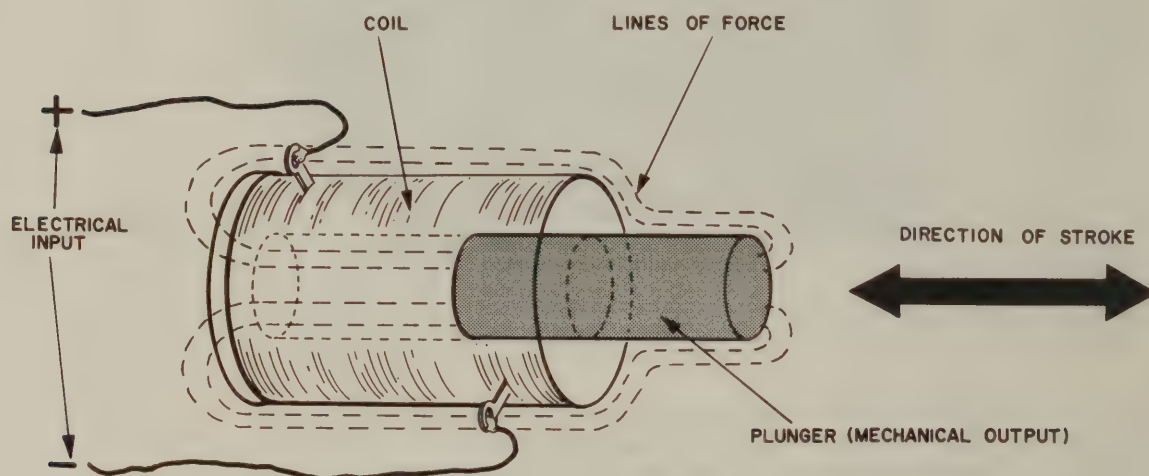


Figure 30

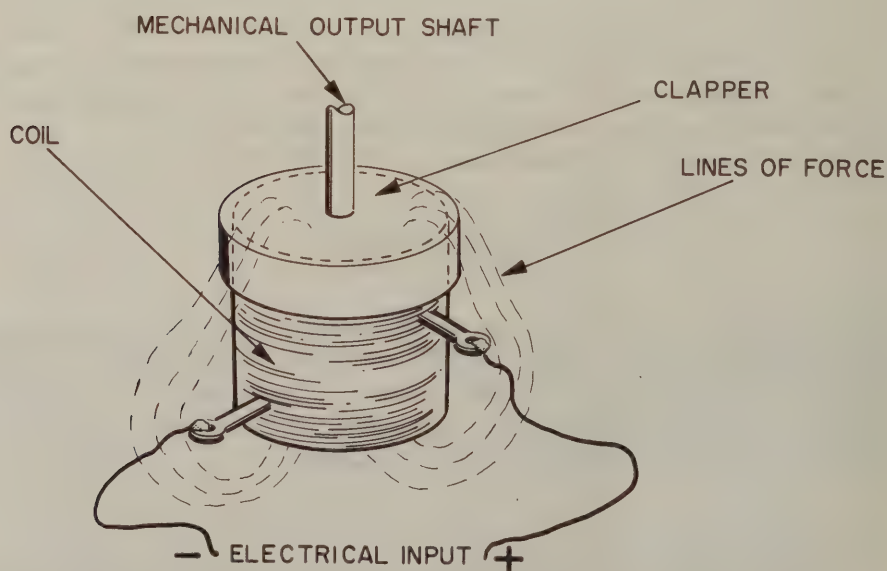
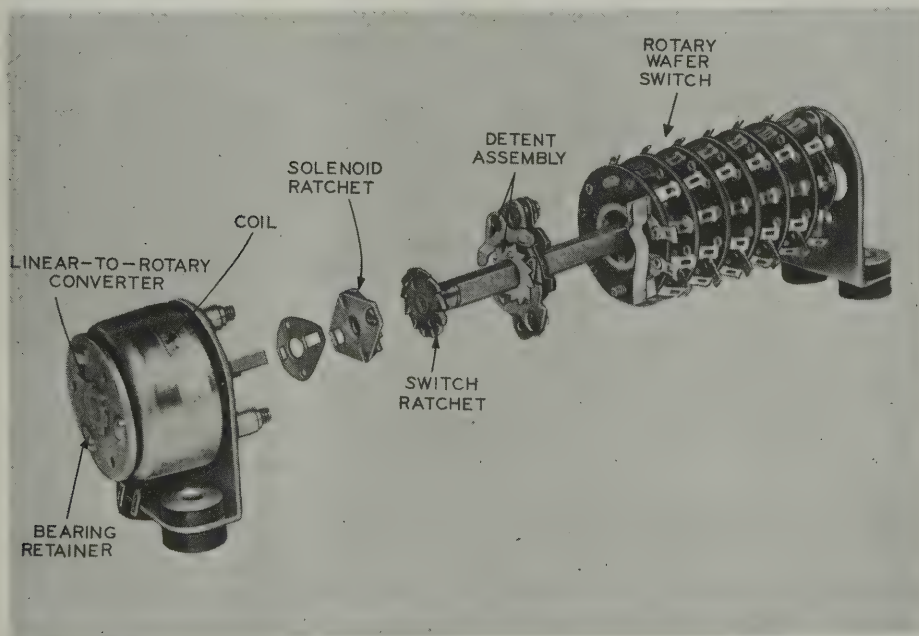
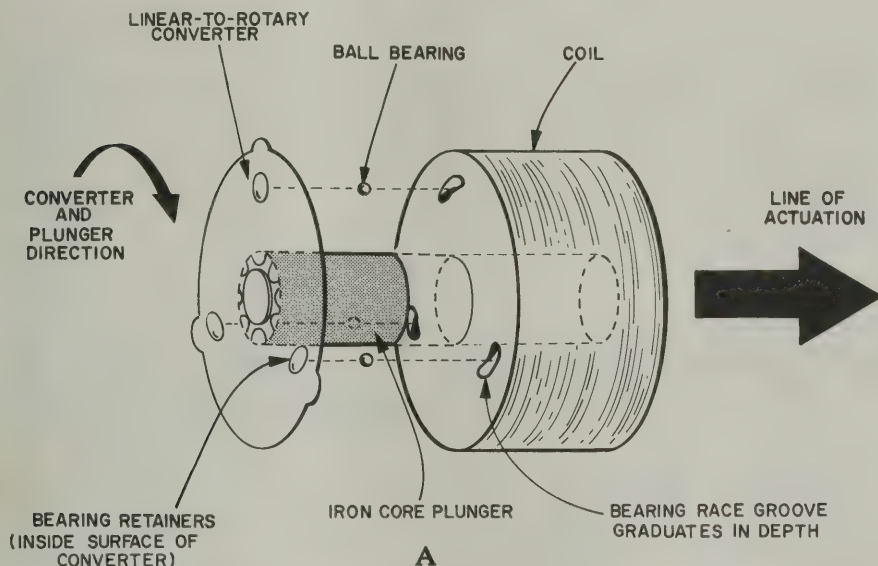


Figure 31

There are three major types of solenoids: (1) **PLUNGER**, (2) **CLAPPER**, and (3) **ROTARY**. The plunger type was just discussed and shown in Figure 30. A straight iron plunger or movable core is drawn into the magnetic circuit. In this way the mechanical motion produced by the solenoid is parallel to the long axis of the coil, and is linear. A plunger solenoid is capable of producing mechanical motions of relatively long distances, called the **STROKE**. Thus, a plunger solenoid produces a linear stroke.



B
Figure 32

Clapper solenoids also produce a linear stroke, but the physical configuration of a clapper solenoid differs from a plunger solenoid, as shown in Figure 30. A clapper solenoid uses a hinged or spring-loaded armature located outside of the coil. Often the armature resembles a "cap" over one end of the coil (Figure 31). The cap may also extend part of the way down the sides of the coil. The armature usually has only a small air gap between it and the coil, since the armature uses a small portion of the magnetic circuit. Since this type of solenoid may have only a small air gap, the stroke is quite short.

The plunger or clapper in a linear solenoid is an ACTUATOR since this is the element which operates external mechanical components. Linear solenoids such as the plunger and clapper types are used to do work through the mechanical movement of the actuator. Extended rods or fasteners are attached to the actuator so that the motion produced by the solenoid can be used to do work.

A rotary solenoid produces a rotational (angular) stroke. There are several ways to accomplish this, but the example shown in Figure 32 illustrates one of the common methods. This device is basically an ordinary plunger solenoid. However, a mechanical LINEAR-TO-ROTARY CONVERTER disc is attached to one end of the plunger, and extends through the coil case. This is shown in the exploded view of the device (Figure 32A). When the converter disc is in the normal position, the ball bearings are held in the bearing race grooves. This permits the disc to "roll", or rotate, a few degrees on the back of the coil case. These grooves, however, are the key to the angular motion obtained from this device. The grooves are not the same depth from one end to the other. They are, instead, INCLINED in ramp fashion. This means that when the coil is energized and the plunger is drawn inward, the bearings roll "downhill" in the grooves. This allows the plunger to be pulled farther into the coil. As the plunger continues in the LINE OF ACTUATION (see Figure 32A), the rolling action of the bearings and disc also causes the plunger to rotate. Thus, the plunger has both a short linear stroke, and a small angular stroke. Generally, however, the linear motion can be absorbed in some way, and the angular motion used to provide the desired turning force.

Figure 32B shows a typical application for a rotary solenoid of this type. In the exploded view (B), a ratchet-type mechanism couples the solenoid actuator and a rotary wafer switch. This device operates in a way quite similar to the rotary stepping switch discussed previously. In fact, the combination presented here is indeed another type of rotary stepping switch. The ratchet mechanism permits continuous 360° rotation of the switch, one step at a time, while the solenoid itself rotates only a few degrees. This is because the ratchet imparts a force on the switch shaft only in one direction. When the solenoid energizes, the slight linear (forward) motion of the actuator engages the ratchet. The clockwise turning (angular) motion then turns the wafer switch one step, or switch position. When current is removed from the

solenoid, the actuator restores by spring action. The actuator then rotates counterclockwise to its de-energized position. The wafer switch does not turn at this time since the ratchet merely "rides over" the teeth. The switch can then be advanced another step by again energizing the solenoid.

ELECTRIC CLUTCHES AND BRAKES

A **BRAKE**, as shown in Figure 33, is a device that exerts a force to retard or stop a rotating member upon application of controlling power. A **CLUTCH**, as shown in Figure 34, is a device which couples two rotating or static members together upon the application of controlling power. The functions of both clutch and brake can be combined to perform both actions within the same unit. This is a **BRAKE-CLUTCH** as shown in Figure 35. A brake-clutch will either disengage a driving force and brake the driven member, or disengage the brake member and apply a driving force. This action can occur simultaneously or within a short time interval. Clutches and brakes may be actuated by an electromagnetic coil and armature, and may therefore be electromechanical devices. The three basic types of electromechanical clutches and brakes are: (1) **FRICITION DISC**, (2) **HYSTERESIS**, and (3) **MAGNETIC PARTICLE**.

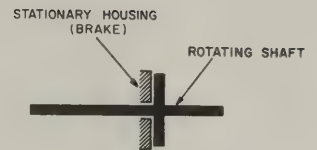


Figure 33

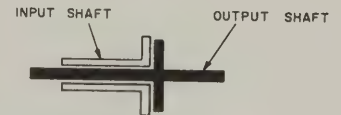


Figure 34

The friction disc clutch and brake is an on-off device. That is, the mechanism is not intended to provide fractional accelerations or continuous slip. As shown in Figure 36A, the friction disc clutch consists of two rotating members. The input shaft is shown by the non-shaded area. The output shaft is the bar-like darkly-shaded area. The entire assembly resembles the basic clutch mechanism described previously in Figure 34. An energizing coil is positioned around the outer shaft as with the other types of clutches. In the friction disc mechanism, an armature is coupled to the inner shaft. When the coil is energized, the armature is pulled toward the input shaft and engages the friction surface between the shafts. Thus the friction disc clutch is similar to an armature relay. By energizing and deenergizing the coil, the air gap between the friction surfaces is controlled and the shafts are coupled or decoupled. The friction disc clutch is the simplest device of the three basic types of clutch. Advantages of the friction disc clutch are fast response time and high speed operation. The friction disc clutch also has higher torque capabilities than either the hysteresis or the magnetic particle mechanism. The same principle applies to the friction disc brake and brake-clutch shown in Figures 36B and 36C.

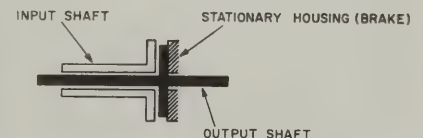


Figure 35

A hysteresis brake or clutch depends upon the magnetic **HYSTERESIS EFFECT** for its operation. When a material such as iron is placed in a magnetic field, a certain amount of energy is required to magnetize the iron. If the field is a rapidly alternating one, the material may become noticeably warm. This phenomenon is similar to mechanical friction, and opposes free motion by the iron member. This causes hysteresis heating due to the expenditure of electromagnetic energy, and thus the dissipation of heat must be

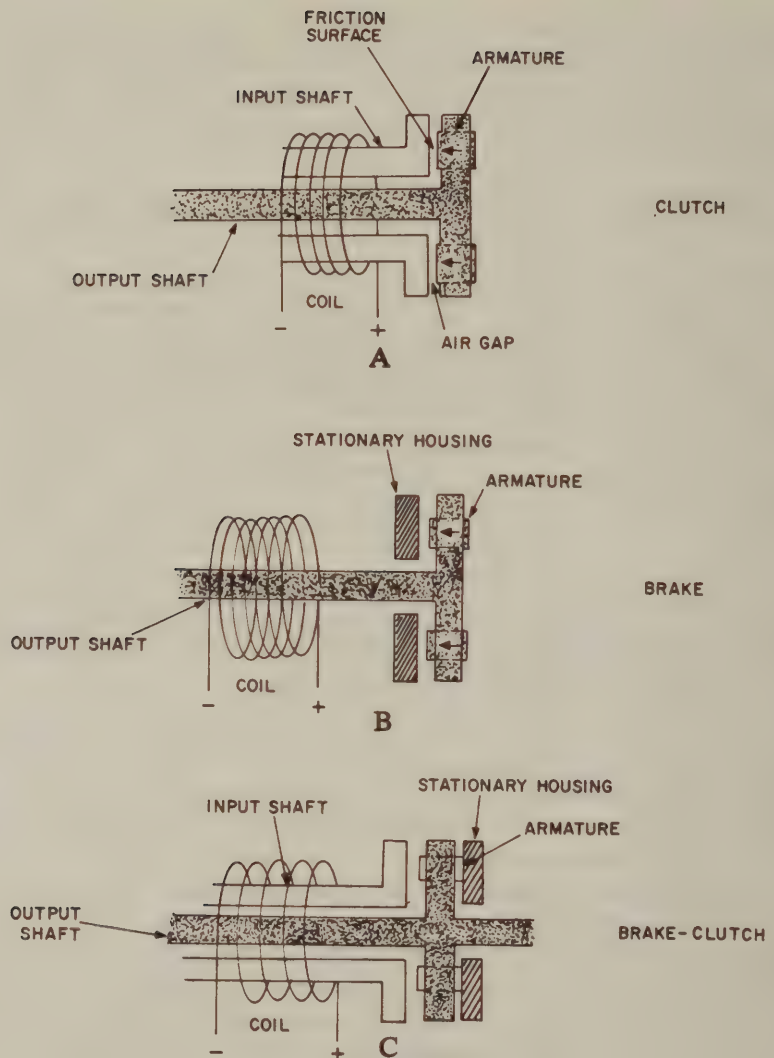


Figure 36

considered in using hysteresis clutches and brakes. A hysteresis clutch or brake is a PROPORTIONAL CONTROL device. That is, the rotating member may be affected to a greater or lesser degree by increasing or decreasing the coil current. The coil current varies the amount of flux induced into the opposite member, and thus the force, or torque exerted upon the opposite member. Provided that the generated heat can be removed, a hysteresis unit can provide slow acceleration of one member, or continuous slip.

A MAGNETIC PARTICLE clutch is shown in Figure 37. The magnetic particle clutch consists basically of a cylindrical plate separated from the rest of the assembly by a fine magnetic particle powder. The plate is part of the inner (output) shaft. The space between the plate and the input shaft is filled with the magnetic powder. The shafts will rotate freely one around the other with the powder between them. A coil, or stator winding, is placed around the shafts.

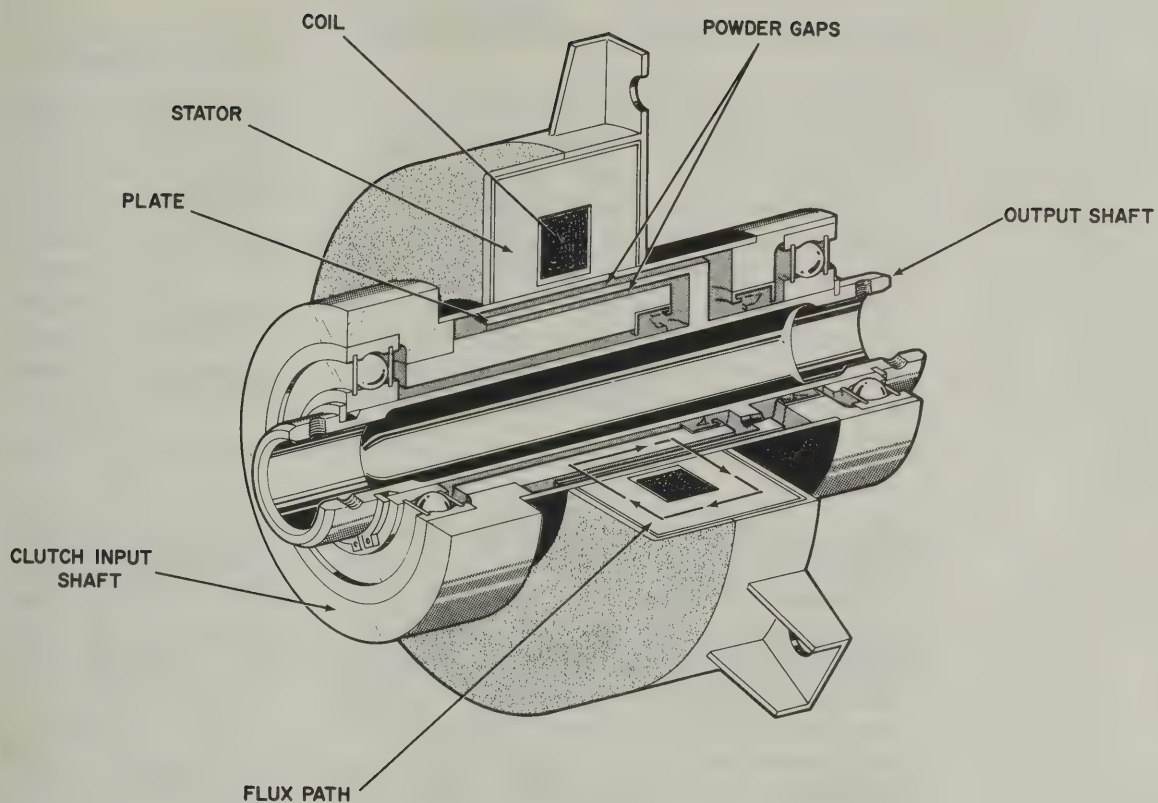


Figure 37

When the coil is energized, the powder becomes partially magnetized and forms a semi-rigid mass. The particle material is then capable of transmitting torque from the moving (outer) shaft to the inner shaft. The iron particle powder has a particle coherence under the influence of a magnetic field which transmits torque almost linearly with the strength of the magnetic field. Thus, an iron particle clutch is a proportional control device as is the hysteresis clutch. The magnetic particle clutch has advantages over the hysteresis clutch in that less heat is generated during operation. As a result, little wear occurs and normal operating life is extended over those of the other types. The particles do become more finely divided over the life of the clutch, but this does not affect the efficiency of the device. Since torque is proportional to the control current, the transmitted torque is proportional to the control current, the transmitted torque is smooth even down to zero, and galling or grabbing cannot occur within the friction surfaces. Magnetic particle brakes work in the same manner, except that one of the shafts is locked in a stationary position. The magnetic powder fills the space between the rotating shaft and stationary housing as before. With no current applied to the stator coil, the shaft rotates freely. When current is applied, the powder becomes semi-rigid and shaft rotation is inhibited. Magnetic particle clutches and brakes have higher torque capabilities than hysteresis devices, but are generally limited to low speed operation.

SUMMARY

The simplest electromechanical device is the manually operated switch. Switches vary widely in type and capabilities. The most basic of these are leaf switches, widely used in the telephone industry. Switches are classed as to function by the number of poles and throws in the switch. A pole is a movable contact. A throw is a stationary contact. Switches are also selected for use by the characteristics of various contact combinations. Common combinations are break-before-make and make-before-break. Some switches are made with internal mechanisms that move only within close tolerances and some which respond to very low pressure. These are precision snap-acting switches. Microswitches are often of this type and may be used to control machinery.

Relays are electromechanical devices by which a variation in the conditions of one circuit affects the operation of other devices in other circuits. Relays all operate on the same basic principles of electromagnetism; however, relays vary as required to fulfill a variety of functions. In addition to the many types of general-purpose relays, some relays are specialized. Telephone relays are a special type which is suitable for switching voice-grade circuits. Telephone systems also use time delay relays extensively, and the technique of using collars and sleeves to achieve the time delays is economical. Other specialized relays are the resonant reed relay and the magnetic reed relay.

Rotary stepping switches are large multiple contact switches which may have hundreds of throws. The actuating mechanism is called a motor magnet, and it can step the switching mechanism across the throws at high speed. The device can also be used to select and program a variety of switching functions simultaneously.

Solenoids are electromechanical devices which qualify as true transducers. Solenoids are transducers which convert electrical energy to mechanical energy. This is, however, not always the ultimate purpose of a solenoid. Solenoids are often used to remotely operate other electromechanical devices such as relays and switches. Used in an actual system, the energy conversion property of a solenoid is often incidental since the ultimate object of this use may be to switch or control electrical circuits.

Electric clutches and brakes use electromagnetic methods to couple two rotating shafts (clutch) or halt the rotation of one member by a stationary member (brake). Important types of clutches are hysteresis, magnetic particle, and friction disc devices. A similar braking device exists for each corresponding clutch mechanism.

The following Practice Exercise questions cover the subjects which you have just studied. They are:

RELAYS

RESONANT REED RELAYS

MAGNETIC REED RELAYS

SOLENOIDS

ELECTRIC CLUTCHES AND BRAKES

18. A resonant reed relay acts as a _____ switch.
19. The frequency at which a reed will vibrate when excited by external oscillations is the reed's _____ frequency.
20. The type of permanent magnet operation that gives the most positive magnetic reed relay action for a given magnet is _____.
21. The actuator of a stepping switch is called a _____.
22. Stepping switches that employ the driving force of a spring are DIRECT DRIVE switches. True or False?
23. A solenoid may be properly classified as a transducer. True or False?
24. The movable iron bar placed inside the coil of a solenoid is called a (a) clapper, (b) plunger.
25. Plunger and clapper solenoids produce a _____ stroke.
26. Energy lost from a rapidly fluctuating electromagnetic field in an iron member is called _____ loss.
27. Both the hysteresis and magnetic particle clutches are _____ devices.

IMPORTANT DEFINITIONS

ACTUATOR — An element in an electromechanical device which initiates mechanical movement and operates the device.

ARMATURE — A moving part which actuates contacts in response to a change in coil current.

BRAKE — A device that, upon application of controlling power, exerts a force that will stop or retard a rotating member, or hold a static member against a specific force.

CLUTCH — A device that, upon activation, couples two rotating or static members together.

COLLAR — A highly conductive sleeve placed over the core of an electromagnet to aid in retarding the establishment or decay of flux within the magnetic path.

CONTACT — A current-carrying part of an electromechanical device which engages or disengages to make, break, or transfer electrical circuits.

DROP-OUT — The value of current, voltage or power for which the contacts of a previously energized relay will always assume the deenergized position.

ELECTROMECHANICAL DEVICE — A device which uses electrical or mechanical energy to affect the operation of one or more other devices or circuits.

HEELPIECE — The part of a relay or coil housing located at the opposite end from the armature.

HYSTERESIS EFFECT — The heating effect of a ferromagnetic material due to a rapidly alternating electromagnetic flux.

KNIFE SWITCH — A switch using a bar like, current-carrying actuator.

LEAF SWITCH — A switch using long flexible springs to engage and disengage contacts.

MICROSWITCH — A small precision switch that can be used to precisely control the position or operation of machinery and circuitry.

OVERTRAVEL — The distance that two contacts travel together after the contacts just touch.

PICK-UP — The value of current, voltage, or power for which the contacts of a previously deenergized relay will always assume the energized position.

IMPORTANT DEFINITIONS (Continued)

PLUNGER — The part of a linear solenoid which is drawn into the coil by solenoid action.

POLE — The portion of a switching device which is associated only with one electrically separate circuit.

RELAY — An electromechanical device which is operated by a variation in the conditions of one electric circuit which affects the operation of other devices in the same or other electric circuits by opening and closing contacts.

RESIDUAL GAP — The length of the magnetic air gap between the pole-face center and the nearest point on the armature when the armature is in the energized position.

SOLENOID — An electromechanical device which converts electrical energy to mechanical energy.

STICKING — The condition of a relay armature that occurs when the relay core has become partially magnetized and will hold the armature operated after coil current has been removed.

THROW — The number of closed switch positions (completed circuits) associated with one electrical circuit in a switching device.

TRANSDUCER — A device which converts energy from one form to another.

PRACTICE EXERCISE SOLUTIONS

1. knife switches
2. contact
3. armature
4. (b) N.C. (normally closed)
5. False — The minimum number of springs required for a transfer combination is three.
6. (a) on-off switching.
7. machinery.
8. 12 positions. Since there is only one movable contact, it is likely that the switch would be used to transfer a voltage or current to 12 separate circuits.
9. (a) armature.
10. sticking
11. pick-up
12. drop-out
13. (a) energizing circuit.
14. (b) heavy duty.
15. power
16. (c) both slow operating and slow releasing.
17. True
18. frequency sensitive
19. resonant
20. perpendicular actuation
21. motor magnet
22. False — Stepping switches that employ the driving force of a spring are called INDIRECT DRIVE switches.

PRACTICE EXERCISE SOLUTIONS (Continued)

- 23. True
- 24. (b) plunger.
- 25. linear
- 26. hysteresis
- 27. proportional control



EXAMINATION
CHECK SHEET

III5A

1. C - IF A SWITCHING DEVICE CONTAINS ONE ELECTRICAL CIRCUIT, IT IS SAID TO HAVE A SINGLE -- POLE.

Another term, throw, is often applied to switching devices to classify the number of closed switch positions. It should not be confused with the term pole, which involves the number of electrical circuits in the device.

2. B - THE SWITCH SHOWN IN FIGURE 5 IS -- DPDT.

The switch has two armature springs, which are movable, and the contacts affixed to them make or break the stationary contacts either above or below.

3. B - STICKING IN A RELAY IS CAUSED BY -- RESIDUAL MAGNETISM.

A relay coil becomes partially magnetized through repeated use of the relay, and the armature, which comes in contact with the core pole face may be held there magnetically after deenergization.

4. C - THE PICK-UP VALUE OF A RELAY IS ALWAYS -- HIGHER THAN THE DROP-OUT VALUE.
The magnetic air gap of an unoperated relay is large, and considerable flux is required to pull the armature toward the coil core.

1. A
B
C
D

5. C - A MAJOR DIFFERENCE BETWEEN TELEPHONE RELAYS AND OTHER TYPES OF RELAYS IS THAT -- TELEPHONE RELAYS SWITCH VOICE-GRADE CIRCUITS.

Telephone relays perform a different function from most relays because voice signals are switched through the contacts rather than power or control signals.

2. A
B
C
D

6. D - RELAYS WITH HEAVY SPRING PILEUPS SHOULD USE -- A SHORT ARMATURE.

The spring force of relays with many springs in the pileup is considerable, and opposes the force of the armature in the energized state. A short armature, placed near the ends of the springs, gives the armature a mechanical advantage in maintaining energization with smaller energizing currents.

3. A
B
C
D

7. A - A SLEEVED RELAY -- IS MUCH SLOWER IN RELEASING THAN A STANDARD RELAY.

A highly conductive sleeve on a relay retards the decay of the electromagnetic field, and delays the relay release.

4. A
B
C
D

8. A - THE ACTUATOR OF A STEPPING SWITCH IS CALLED THE -- MOTOR MAGNET.

A stepping switch uses an electromagnetic coil to drive the stepping mechanism, sometimes at high speed, and is analogous to a small motor.

5. A
B
C
D

9. B - THE ACTUATOR OF A LINEAR SOLENOID IS THE -- PLUNGER.

The plunger is the iron bar which is drawn into the coil upon energization. This mechanical action is utilized in doing work.

6. A
B
C
D

10. C - A MAGNETIC PARTICLE CLUTCH IS A PROPORTIONAL CONTROL DEVICE BECAUSE -- THE MAGNETIC PARTICLE MATERIAL BECOMES MORE RIGID AS THE COIL CURRENT INCREASES.

The magnetic particle material has a particle coherence which transmits torque almost linearly with the strength of the magnetic field.

7. A
B
C
D

8. A
B
C
D

9. A
B
C
D

10.

PRACTICE EXERCISE SOLUTIONS (Continued)

- 23. True**
- 24. (b) plunger.**
- 25. linear**
- 26. hysteresis**
- 27. proportional control**

QUESTIONS

IMPORTANT — These instructions **MUST** be accurately followed to avoid loss, or errors in grading.

Indicate your answer on this sheet by filling in the box for the most correct answer to each question, as illustrated in the example below.

When all questions have been answered, place the answer card in the proper position to line up the boxes on the card with the boxes on the sheet.

Next, copy the complete lesson code into the space provided on the card, and fill in the answer boxes to correspond with those previously filled in on this sheet.

Before mailing, be certain your correct student number, name and address appear on the card.

LESSON CODE
1115A

Example: A fish commonly found in home aquariums is the
 A ☒ B ☐ C ☐ D ☐
 A. gold fish. B. shark. C. barracuda. D. whale.

1. A ☒ B ☐ C ☐ D ☐ **If a switching device contains one electrical circuit, it is said to have a single**
 (A) throw. (B) armature. (C) pole. (D) N.C. contact.

2. A ☐ B ☒ C ☐ D ☐ **The switch shown in Figure 5 is**
 (A) SPST. (B) DPDT. (C) SPDT. (D) DPST.

3. A ☐ B ☒ C ☐ D ☐ **Sticking in a relay is caused by**
 (A) spring action. (B) residual magnetism. (C) the heelpiece. (D) the magnetic air gap.

4. A ☐ B ☐ C ☐ D ☐ **The pick-up value of a relay is always**
 (A) lower than the drop-out value. (B) the same as the drop-out value. (C) higher than the drop-out value. (D) one-half the drop-out value.

5. A ☐ B ☐ C ☒ D ☐ **A major difference between telephone relays and other types of relays is that**
 (A) the telephone relay must be very heavy duty. (B) telephone relays are less efficient than other types of relays. (C) telephone relays switch voice-grade circuits. (D) telephone relays switch very high currents.

6. A ☐ B ☐ C ☐ D ☐ **Relays with heavy spring pileups should use**
 (A) a long coil. (B) a long armature. (C) a collar. (D) a short armature.

7. A ☒ B ☐ C ☐ D ☐ **A sleeved relay**
 (A) is much slower in releasing than a standard relay. (B) has a higher pick-up value than a standard relay. (C) has two coils. (D) is much slower in operating than a standard relay.

8. A ☒ B ☐ C ☐ D ☐ **The actuator of a stepping switch is called the**
 (A) motor magnet. (B) ratchet. (C) driving spring. (D) pawl.

9. A ☐ B ☐ C ☒ D ☐ **The actuator of a linear solenoid is**
 (A) the ratchet. (B) the plunger. (C) the coil. (D) the disc.

10. A ☐ B ☒ C ☐ D ☐ **A magnetic particle clutch is a proportional control device because**
 (A) the device has the highest torque capabilities of all the various types of clutches. (B) the device is used to stop the rotating member smoothly without stalling or grabbing. (C) the magnetic particle material becomes more rigid as the coil current increases. (D) the magnetic particle material becomes proportionally more finely divided as the device is used.

UNIT EXAMINATION FOR COURSE 252

This examination contains a number of different type questions: multiple choice, true or false, fill in the blank and direct question. To assist us in grading please print your answers clearly on the answer sheet. **COMPLETE THE PORTION STATING YOUR NAME, ADDRESS AND STUDENT NUMBER ON THE ANSWER SHEET AND SEND ONLY THE ANSWER SHEET FOR GRADING.** Retain the examination questions for your records.

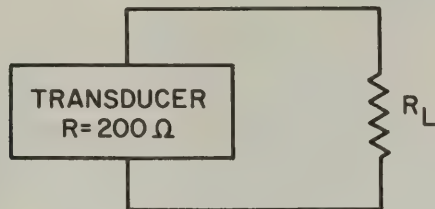


BELL & HOWELL SCHOOLS

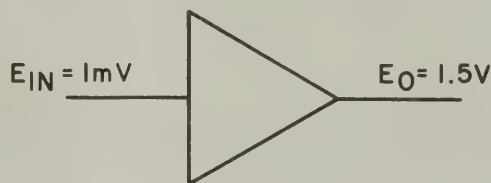
COPYRIGHT 1971
BELL & HOWELL SCHOOLS, INC.



1. To transfer maximum power from the transducer to the load, R_L , shown in the circuit below, what is the required value of R_L ? _____

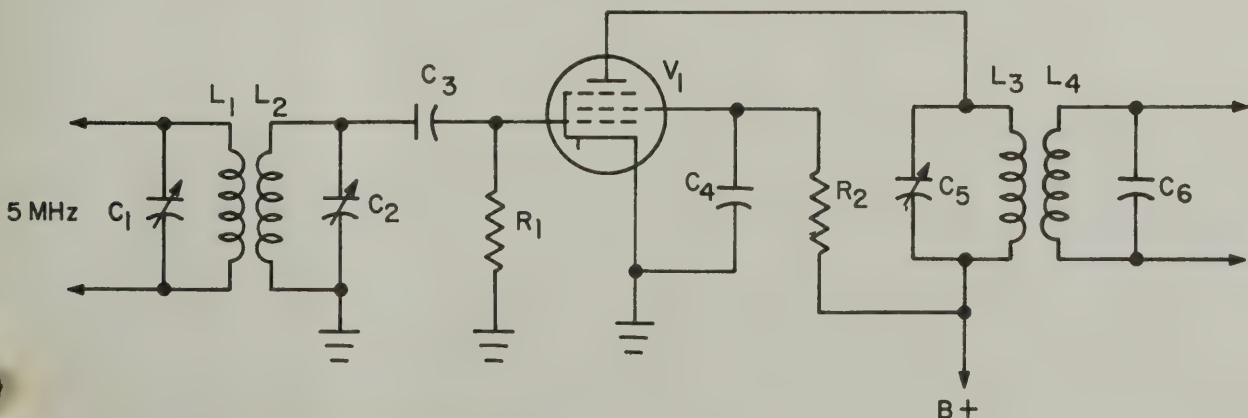


2. What is the voltage gain in the amplifier shown in the circuit below?

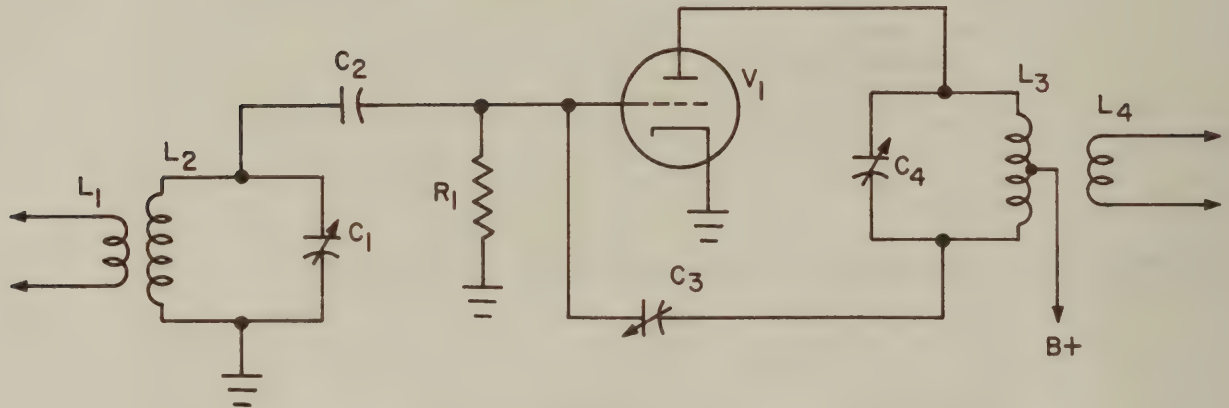


3. In a broadcast television system, the picture carrier is frequency modulated. True or False?
4. When the coupling is reduced below the critical value in a double-tuned i-f amplifier, does the bandwidth increase or decrease? decrease
5. What is the efficiency of a power amplifier if the input power is 1 watt, the output power is 40 watts, $E_b = 650$ volts and $I_b = 130$ mA?

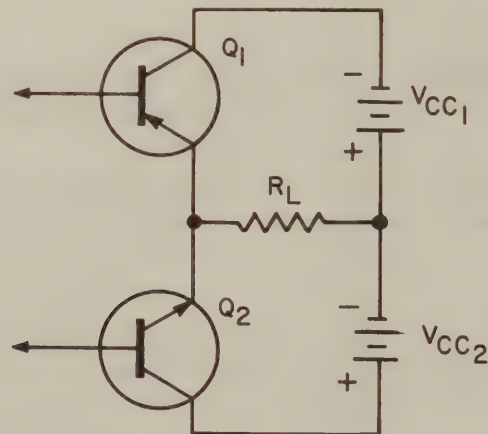
6. For the circuit shown below to operate as a frequency tripler, to what frequency should the C_5 and L_3 tank circuit be tuned? 15 MHz



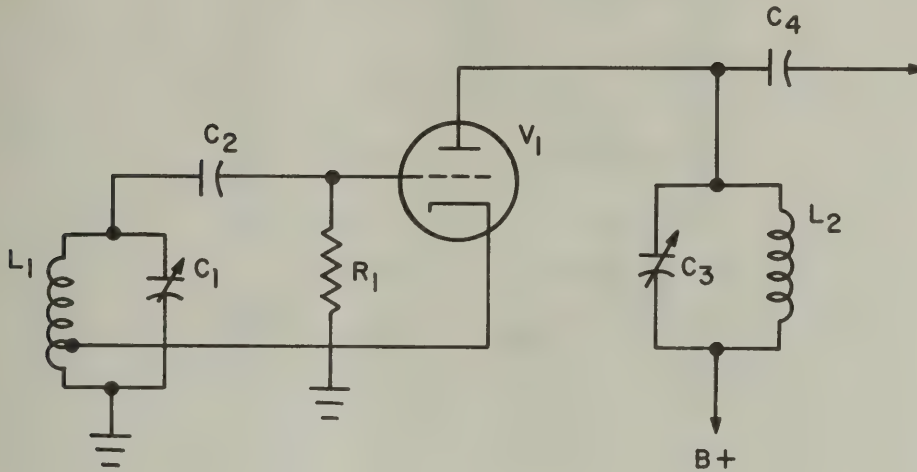
7. In the circuit shown below
 (A) grid neutralization is employed. (B) Rice neutralization is employed.
 (C) Hazeltine neutralization is employed. (D) C_2 serves as the neutralizing capacitor.



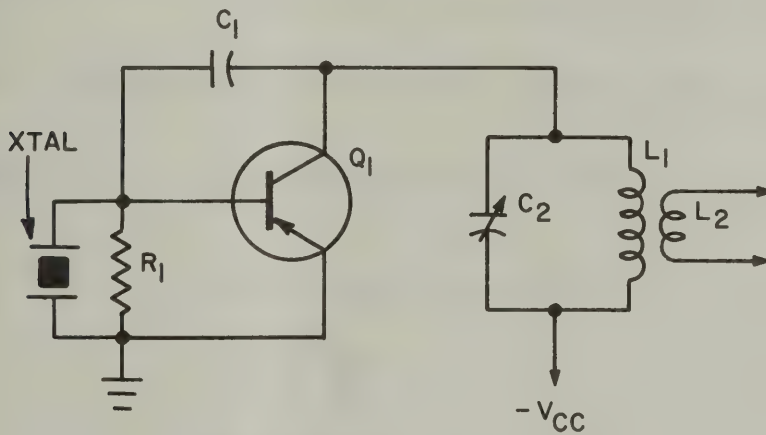
8. To obtain good efficiency a transistor push-pull audio amplifier would generally be operated Class AB or B. True or False? True
9. Does the transistor push-pull amplifier configuration shown below require a phase inverter? No



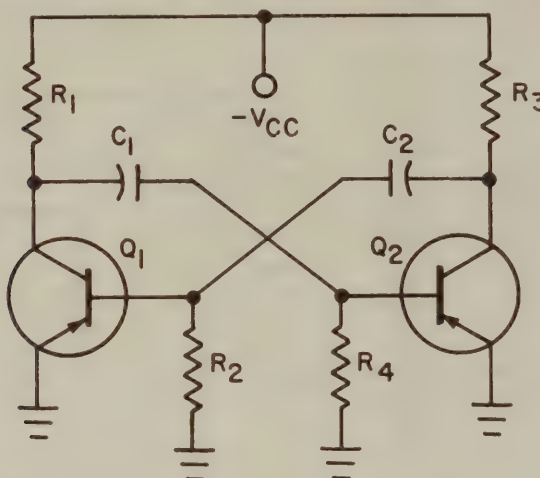
10. The oscillator circuit shown below is a
 (A) series fed Hartley. (B) tuned-plate tuned-grid. (C) shunt fed Hartley. (D) shunt fed Colpitts.



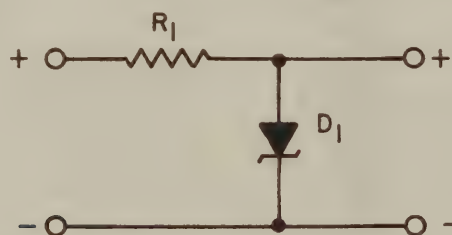
11. The crystal oscillator circuit shown below is similar to a
 (A) tuned-plate tuned-grid. (B) electron coupled. (C) series fed Hartley. (D) shunt fed Colpitts.



12. In the circuit shown below, will increasing the value of C_1 and C_2 increase or decrease the operating frequency? decrease



13. For good frequency stability in an LC oscillator, should the L to C ratio be high or low? low
14. In a series-fed transistor r-f oscillator, collector current flows through all or part of the tank circuit. (True) or False?
15. In a junction type field effect transistor, is the gate forward or reverse biased? _____
16. Is the zener diode, D_1 , properly connected in the circuit shown below? _____



17. A silicon controlled rectifier, SCR, can be turned on or off by the signal applied to the gate terminal. (True) or False?
18. To increase the output of a separately excited dc generator, would you increase or decrease the field current? increase
19. Does the speed of a shunt type dc motor increase or decrease if the field voltage is reduced by adjusting a speed control? increase

20. Which of the following integrated circuit logic systems has the poorest noise margin?
(A) T²L. (B) DTL. (C) RTL. (D) TTL.

